

Convolutional Neural networks



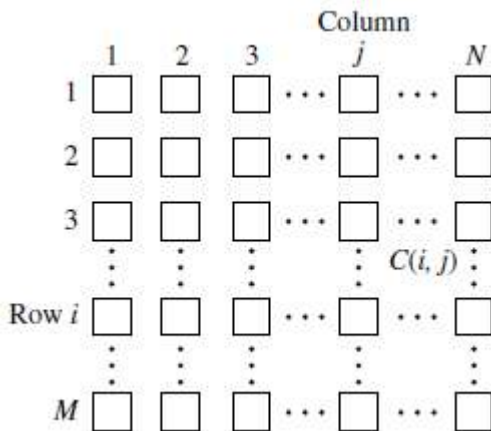
András Horváth

Mire jó egy CNN architektúra?

Turing teljes, de mit old meg hatékonyan?

2D $\rightarrow$ 2D reprezentáció. Térbeli és téridőbeli lenyomatok („feature”)-ök előállítása

A standard CNN architektúra cellák M*N-es reguláris hálózata



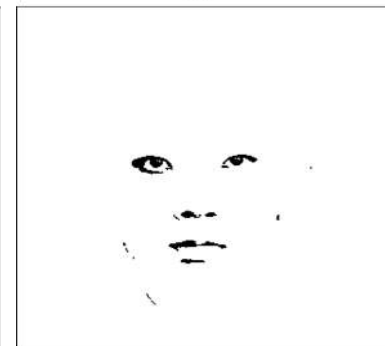
(a)



(a)



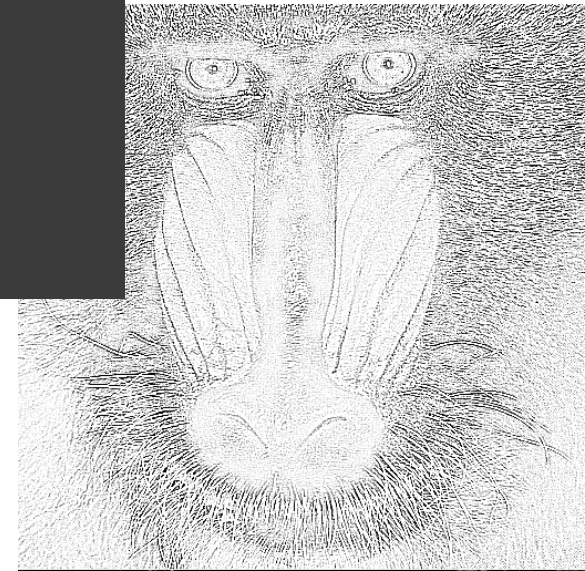
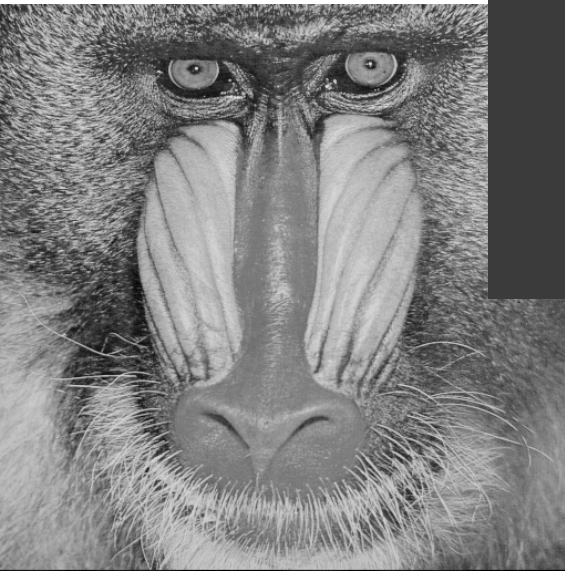
(b)



(b)

Térbeli lenyomat

Gradient



Retina felépítése

Gerinces retina: Inverz retina (ellentétben a fejlábúakkal, külön fejlődési ág) az érzékelő sejtek a leghátsó rétegben találhatóak

Csapok (cones):

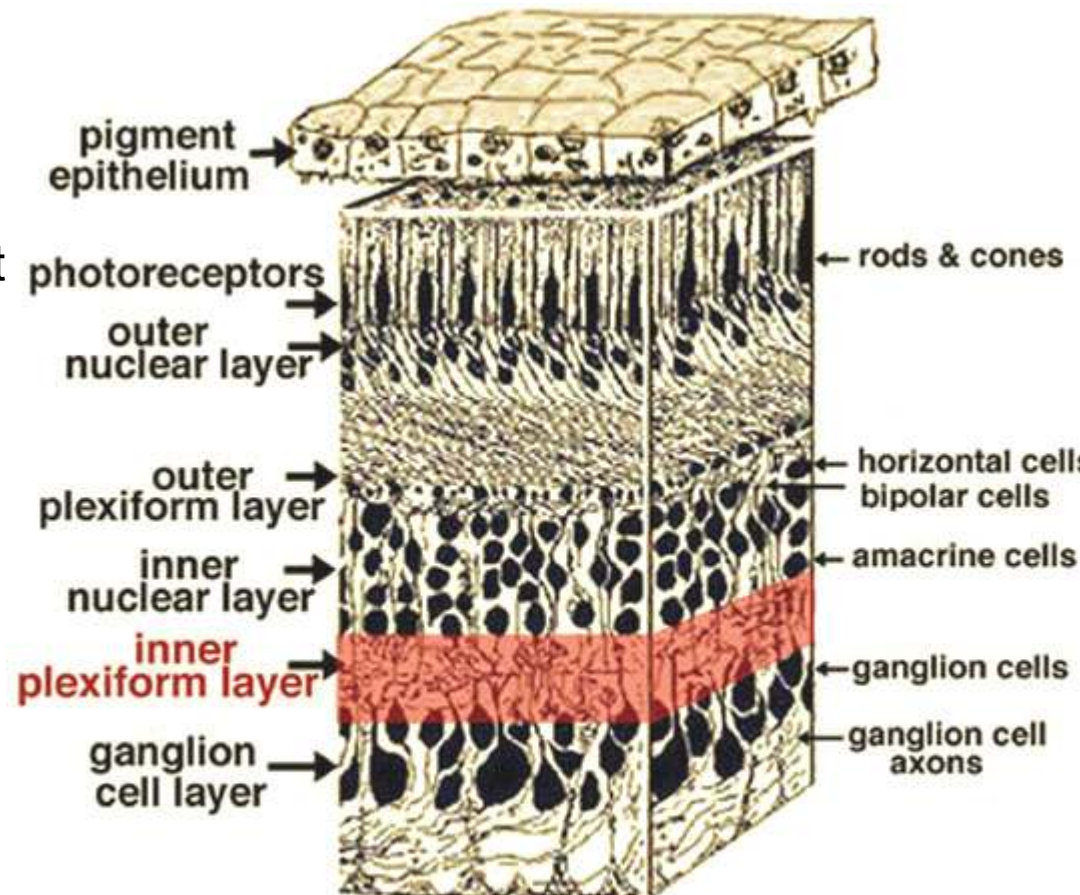
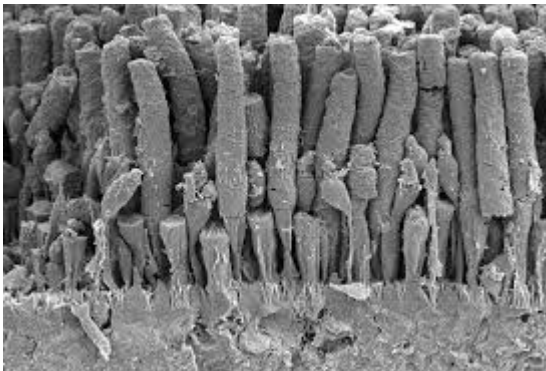
6 millió, Erős megvilágítás mellett

Fovea-ban, színlátás

Pálcikák (rods):

120 millió, gyenge megvilágítás mellett
környéki látás, egy szín

Ganglion, amacrin, bipoláris,
horizontális sejtek



Hierarchikus struktúra a retinában

100x több fotoreceptor, mint ganglion sejt:

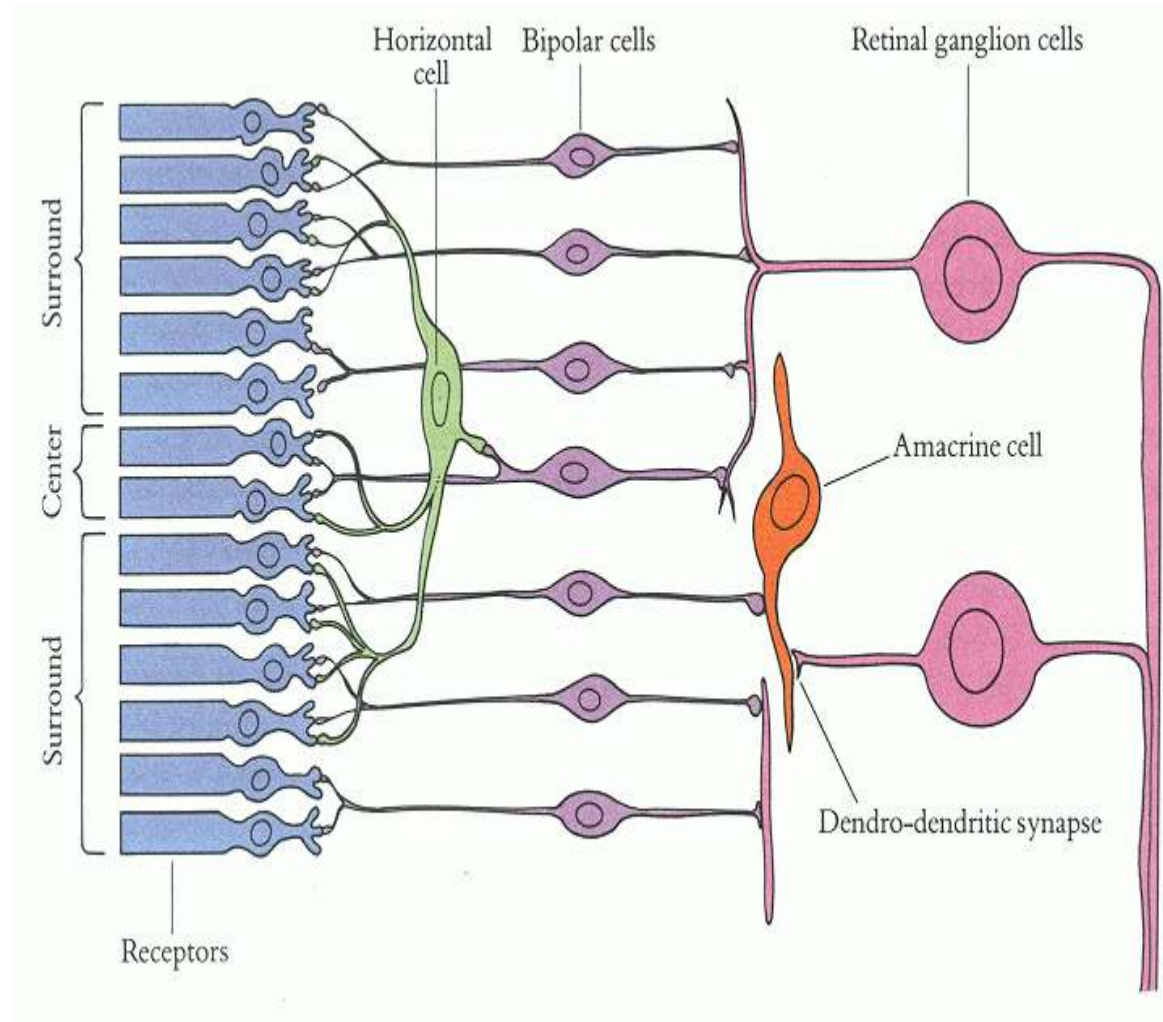
Nagymértékű tömörítésre van szükség

Nem pixel szinten látunk:
Területeket változásokat érzékelünk

Sárgafolt

On/Off

Serkentő, gátló régiók



Mit várunk egy felismerő rendszertől?

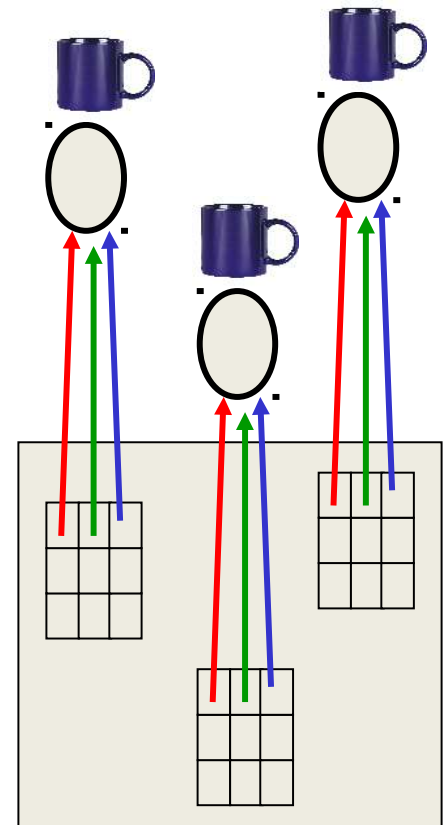
Ismételt lenyomatok (The replicated feature approach)

Nagyon népszerű neurálsi hálózatok esetében

És általában döntési folyamatokban

Ugyanazon lenyomat kinyerő struktúrát használjuk az adaton különböző 'helyzetekben'

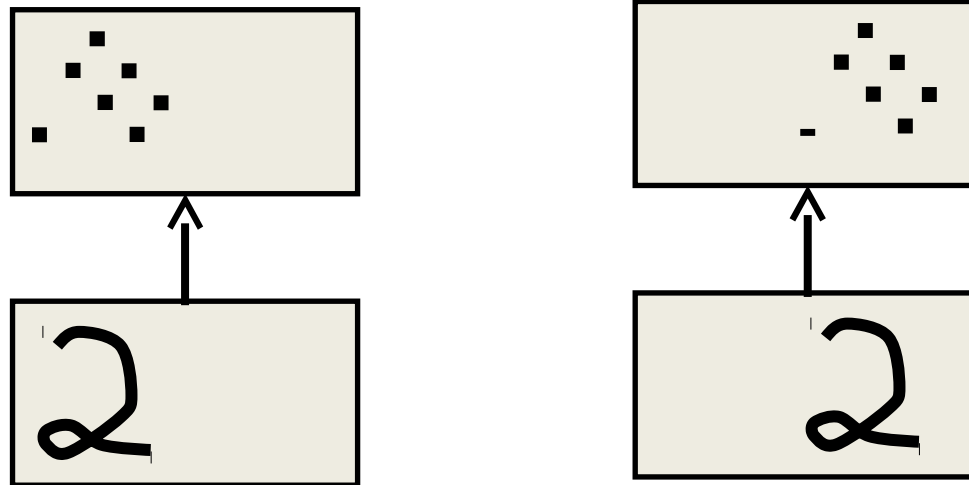
- Különböző pozíciókban azonos választ váránk
- Azonos választ váránk különböző méretű és orientációjú objektumokban (nehéz és számításigényes)
- Nagyon leegyszerűsíti a tanítást, sokkal kevesebb szabad paraméter lesz



Mit érünk el a lenyomat kiszámításának topografikus ismétlésével?

Nagyon leegyszerűsíti a tanítást, sokkal kevesebb szabad paraméter lesz

Invariant knowledge: A tanítás esetén ha egy régióban hasznosnak bizonyult egy feature és ezt a rendszer megtanulta, akkor az összes pozícióban hasznos lesz.



Ha a feature-ök detekcióját homogén CNN struktúrával hajtjuk végre, akkor a detekció biztosan invariáns lesz a translációra

Deep Learning

Általános feature hierarchiákat tanul meg

'Hagyományos tanulás' (Shallow learning) Különböző feature-ök tulajdonságok, melyik hasznosítható?

Kép. Pixel intenzitás, élék, fourier stb - A cél magának a reprezentációnak a megtanulása

Deep learning vs. Shallow learning

- | | |
|---|---|
| <ul style="list-style-type: none">• Focus on entities and connections• Relates previous and new knowledge• Uses reflection to relate theory with experience• Creates understanding, meaning and new ideas• Leads to positive emotions and attitudes about learning and self | <ul style="list-style-type: none">• Focus on unrelated details• Information is simply memorized• Facts and concepts accepted unreflectively• Aims to pass (or perform) instead of understanding• Leads to negative emotions and attitudes about learning and self |
|---|---|

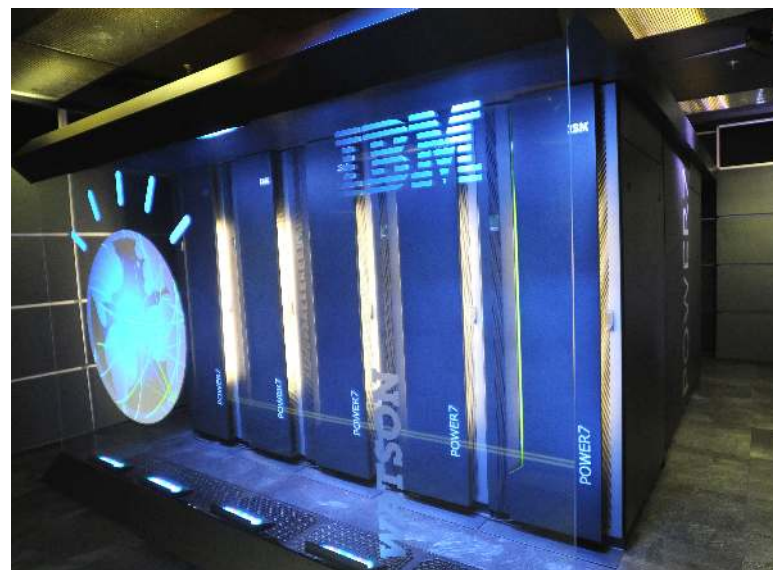
Deep Learning

IBM Watson

A kutatók célja az volt, hogy a természetes szövegértést javítsák a gépek esetében

Neurális hálózatot és deep learning-et használtak egy nagy adathalmazon

A gép megérti a számára feltett kérdéseket és számos esetben jobban teljesít, mint az emberek (Jeopardy)



Deep Learning



Deep Face

Facebook arcfelismerő: DeepFace:
Closing the Gap to Human-Level
Performance in Face Verification

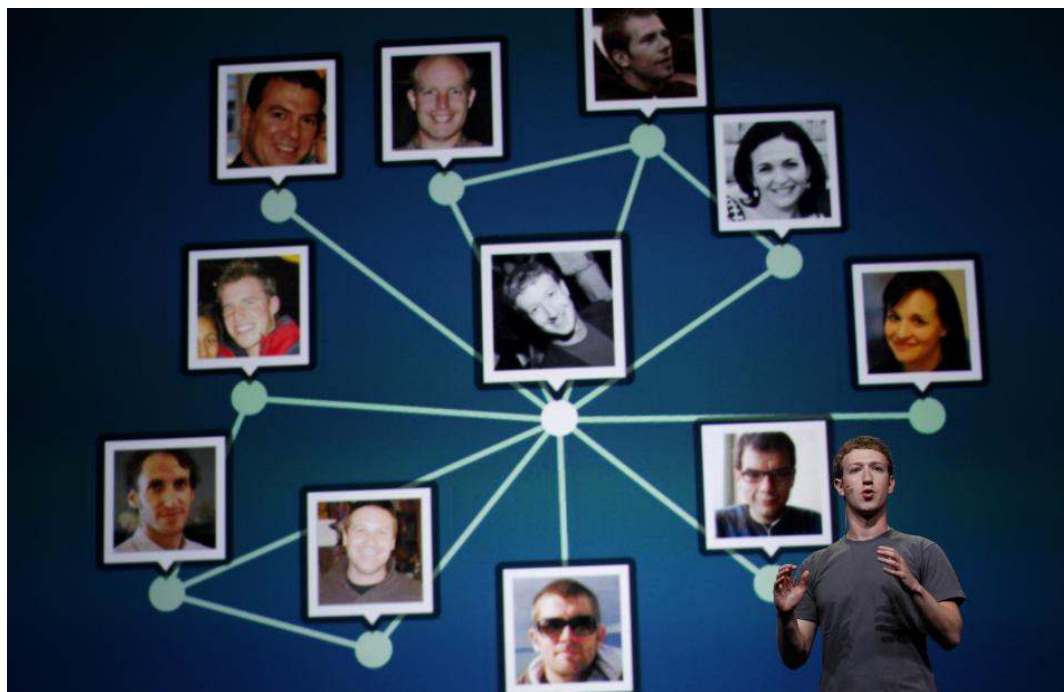
97.35%-os pontosság, egy átlagos ember
98%-os pontossággal dolgozik

Who's in These Photos?

The photos you uploaded were grouped automatically so you can quickly label and notify friends in these pictures. (Friends can always untag themselves.)



Arcfelismerés Shallow learning használatával



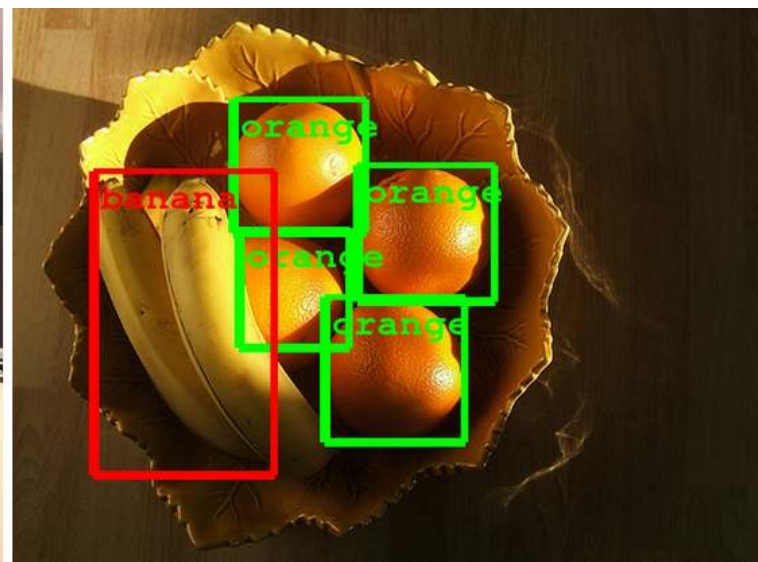
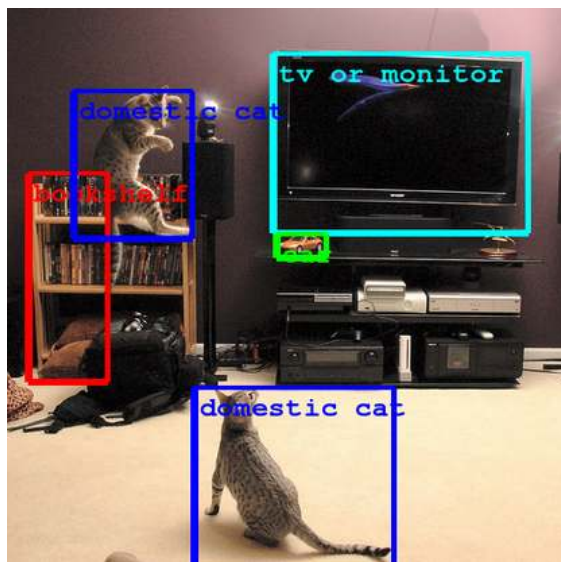
GoogLeNet

Általános objektum felsimerő algoritmus

Az ImageNet adathalmazon tesztelve

Több, mint 200 különböző objektum felismerésére tanítva

37%-os pontosság



Általános lenyomatok előállítása

Az adat lokális stuktúráit szeretnénk felismerni
(akár térbeli akár időbeli): képek, beszéd esetében

Mi a legáltalánosabb olyan tulajdonság, ami:

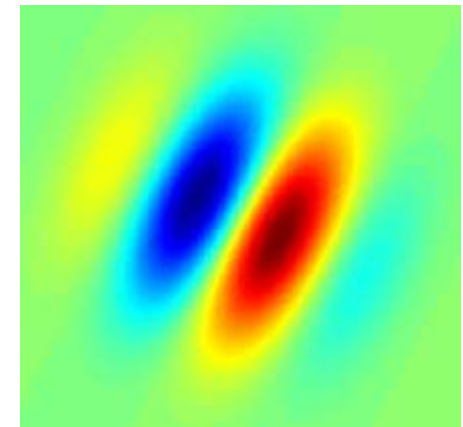
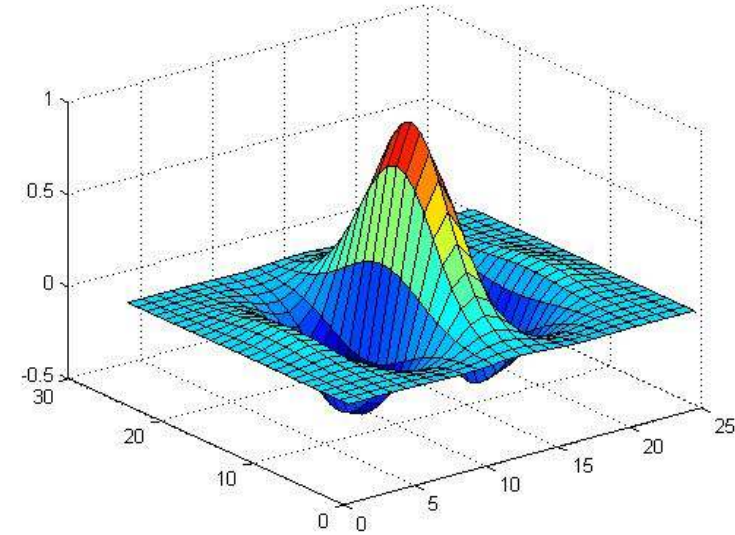
- lényeges a detekció szempontjából
- A legtöbb feature kikeverhető belőle
- Invariáns az értékek nagyságára

Ezek általában egyszerű változásdetektorok

Gábor kernel-ek (Egy egyszerű feed-forward template

- csak B template-tel számolható)

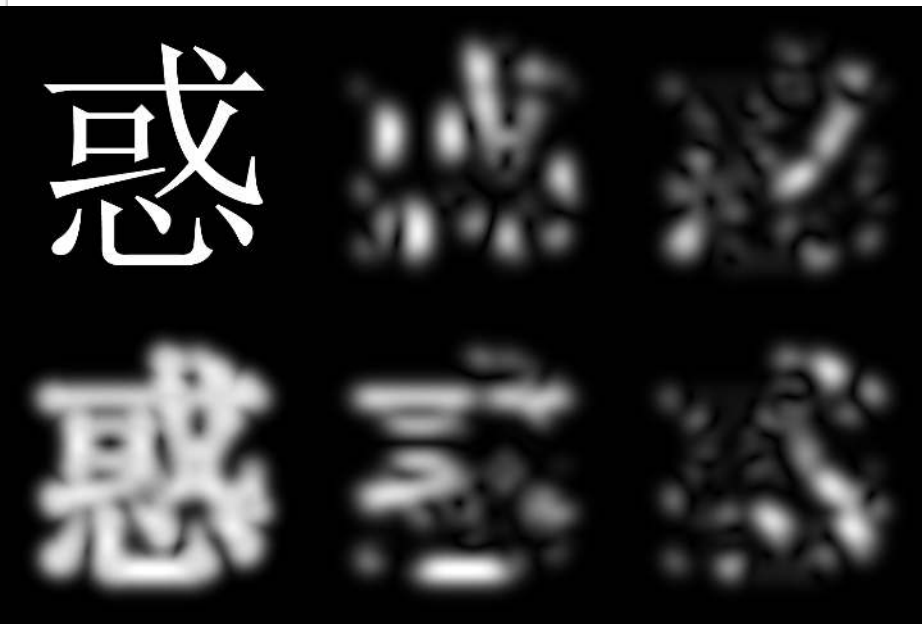
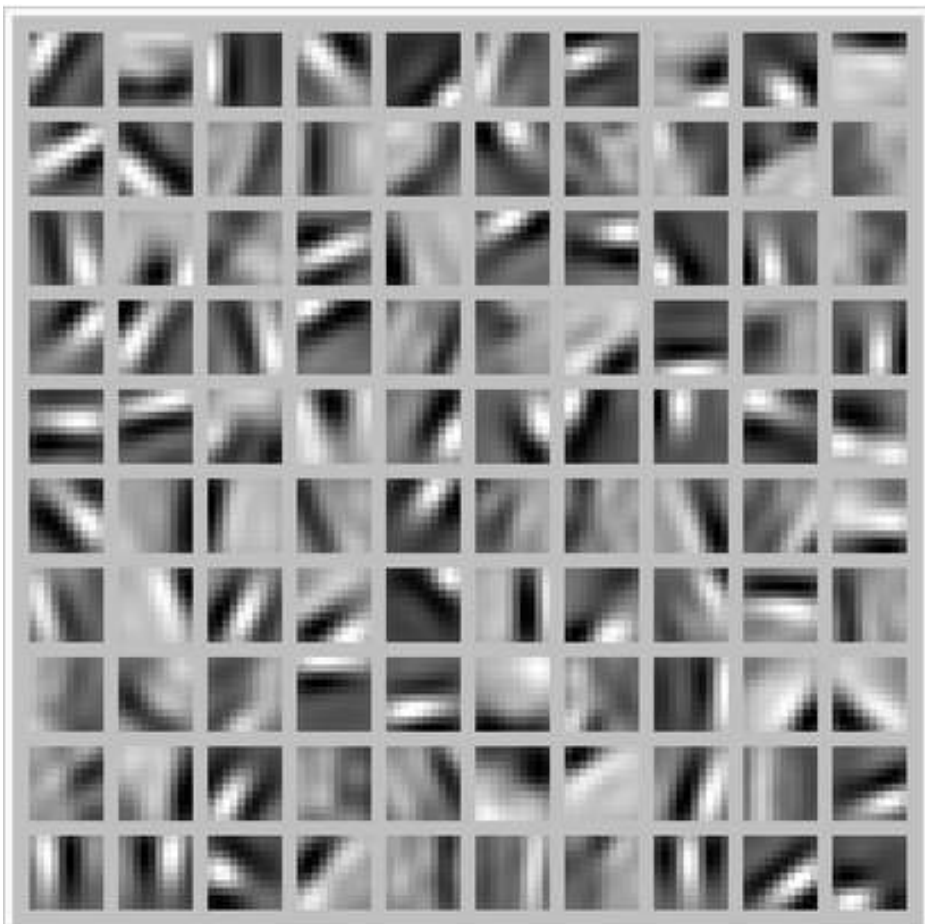
$$g(x, y; \lambda, \theta, \psi, \sigma, \gamma) = \exp\left(-\frac{x'^2 + \gamma^2 y'^2}{2\sigma^2}\right) \exp\left(i\left(2\pi\frac{x'}{\lambda} + \psi\right)\right)$$



Deep Learning

Általános feature hierarchiákat tanul meg

Képek esetében ezek a következőképpen néznek ki:



Neurális hálók használata detekcióban

Kunihiko Fukushima (1980) – Neo-Cognitron



Yann LeCun (1998) – Convolutional Neural Networks

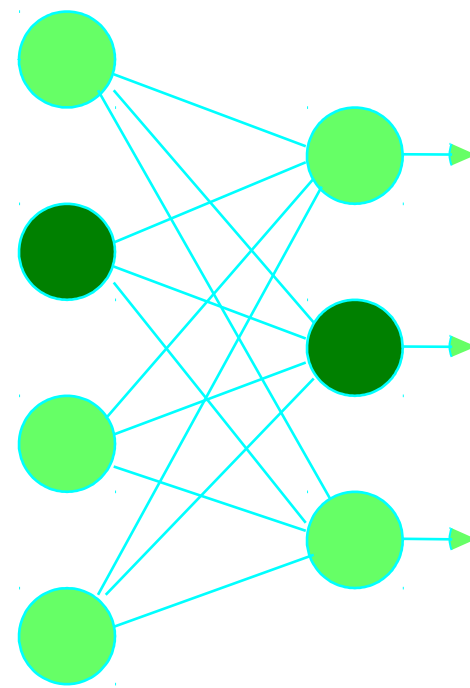
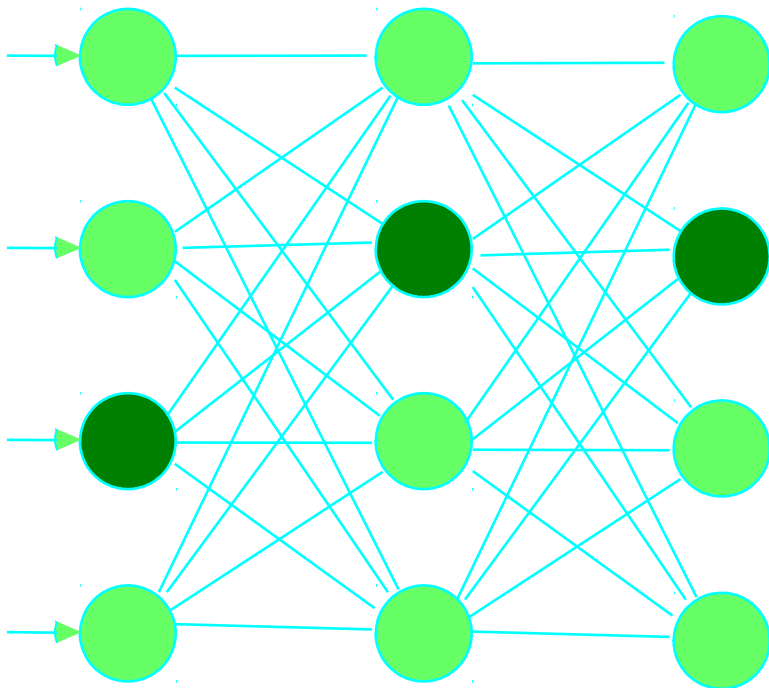
MLP (multilayered perceptron)
backpropagation-nel

- Bízató kezdeti eredmények, de megállt jóval az emberi képességek szintje alatt
 - Rendkívül lassú tanítás (gyorsabb felismerés), bonyolult hálózatok
 -



Hogyan lesz döntés a lokális feature-ökből?

Hagyunk-e egy általános struktúrát, ahol sok paraméterünk lesz és nehéz tanítani, vagy élünk-e megkötésekkel az architektúrára nézve?



Mintavételezés - pooling

Hogyan tömörítsük/alakítsuk döntéssé a térbeli lenyomatot?

Mintavételezés pooling

Két doglot szeretnénk elérni:

- Magasabb szinten a feature-ök magasabb szintjén reprezentálni (pl első szint knotraszt/élek), második szint élek egymáshoz viszonyított elhelyezkedése stb...
 - Kismértékű transzlációra magasabb szinteken is ugyanazt a választ várjuk
- Pooling – egy területen közelítsük a lenyomatot statisztikai leírókkal (a pixelek átlagával vagy maximumával)
- Ez nagymértékben lecsökkenti az elemek számát a következő szinten.
 - Problem: Nagyon sok pooling után elveszítjük a pontos térbeli reprezentációt

Convolutional Neural networks

C (Convolution) és S rétegek (Pooling) váltakozása (Nemlinearitások alkalmazása) a végén egy Multilayer Perceptron a felismeréshez

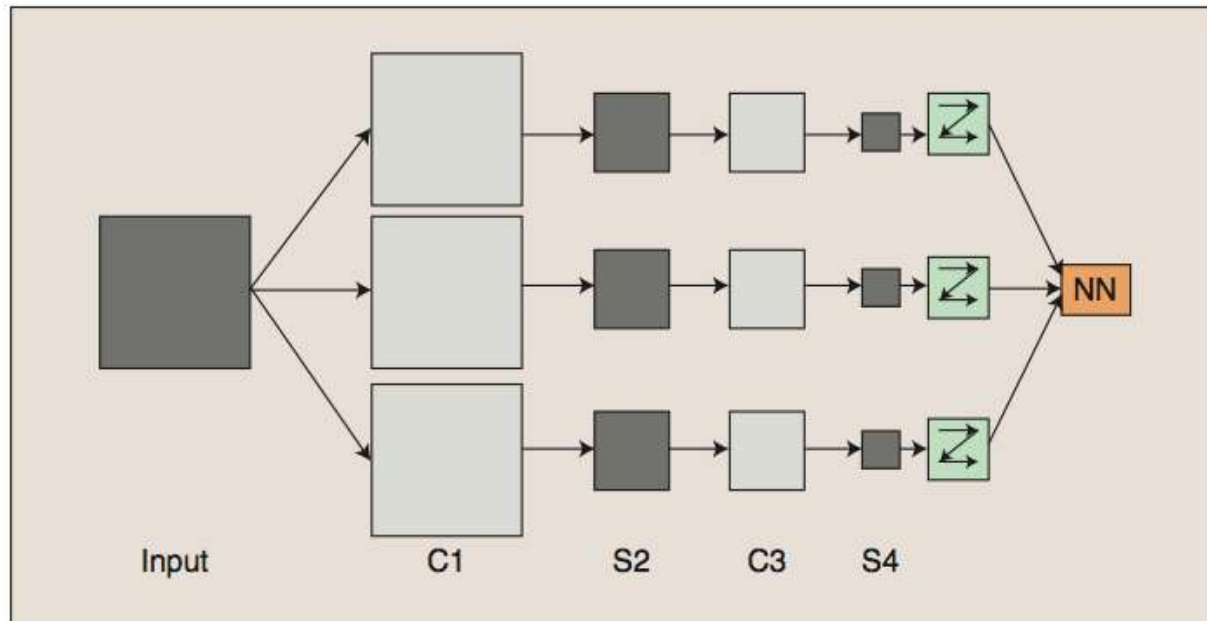
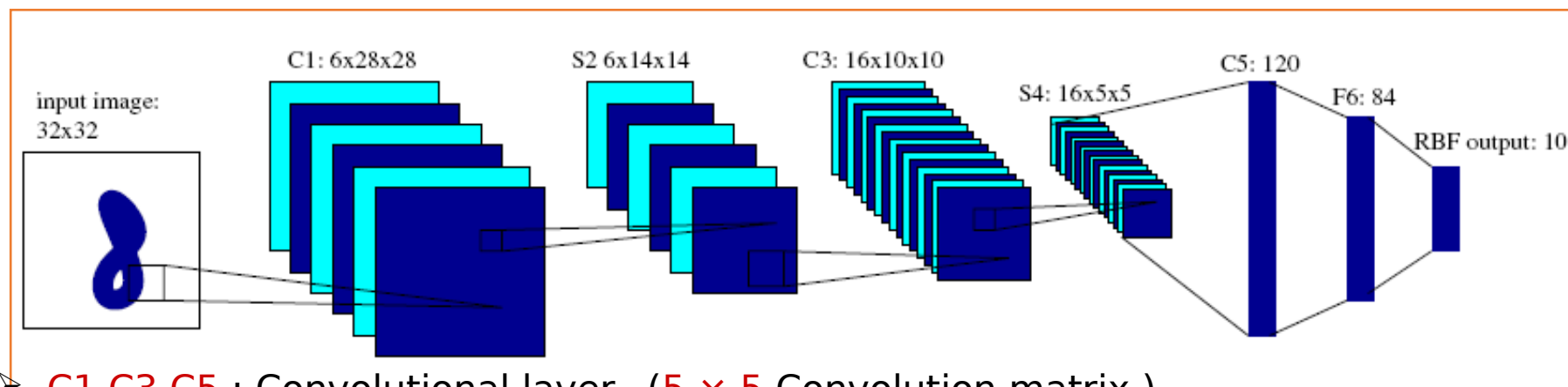


FIGURE 2 Conceptual example of convolutional neural network. The input image is convolved with three trainable filters and biases as in Figure 1 to produce three feature maps at the C1 level. Each group of four pixels in the feature maps are added, weighted, combined with a bias, and passed through a sigmoid function to produce the three feature maps at S2. These are again filtered to produce the C3 level. The hierarchy then produces S4 in a manner analogous to S2. Finally these pixel values are rasterized and presented as a single vector input to the “conventional” neural network at the output.

Convolutional Neural networks

Egy viszonylag sok rétegű hálózat tanítása általános feature-ökkel „shallow learning” esetén összetett feature-ök, de 1-2 rétegű hálózat A későbbi rétegek egyre bonyolultabb feature-öket írnak le:

- Első réteg élék, kontrasztok leírása
- 2nd layer learns higher order features (combinations of first layer features, combinations of edges, etc.)
- A rétegeke felügyeleten tanulással (unsupervised learning) így (az adaathalmaztól függően) általános feature-ök keletkeznek
- Az utolsó réteg tanítása felügyelt, itt megmondjuk, hogy milyen kimenetet szeretnénk egy adott bemenetre



➤ **C1, C3, C5** : Convolutional layer. (**5 × 5** Convolution matrix.)

➤ **S2, S4** : Subsampling layer. (by factor **2**)

➤ **F6** : Fully connected layer.

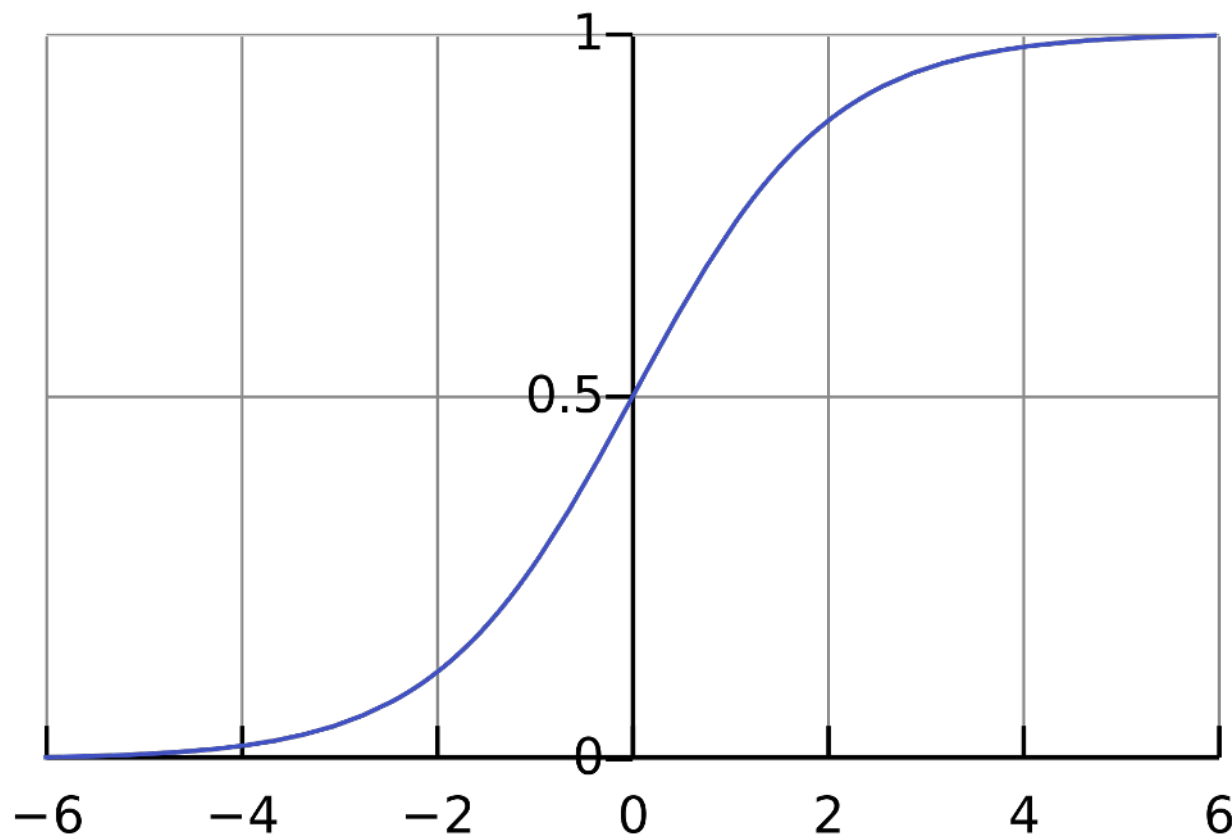
➤ **187,000** összeköttetés

➤ **14,000** állítható súly (10X)

A rétegek között használt nemlinearitás

A rétegek közötti értékek korlátosságának megőrzésére egy a standard CNN dinamikához hasonló nemlinearitást használnak

$$f(x) = \frac{L}{1 + e^{-k(x-x_0)}}$$



Ugyanezt a nemlinearitást használják CNN dinamikákban is (feature kinyerés szempontjából nincs kvalitatív különbség a kettő között)

Hasonlít a biológiai struktúrákhoz

Hierarchikus felépítés az agyban önamguktó lszervező struktúrák

- Topografikus területeken – és feature számítások (retina – V1)
Mintavételek ezen területekről és komplex döntéshozó rendszerek

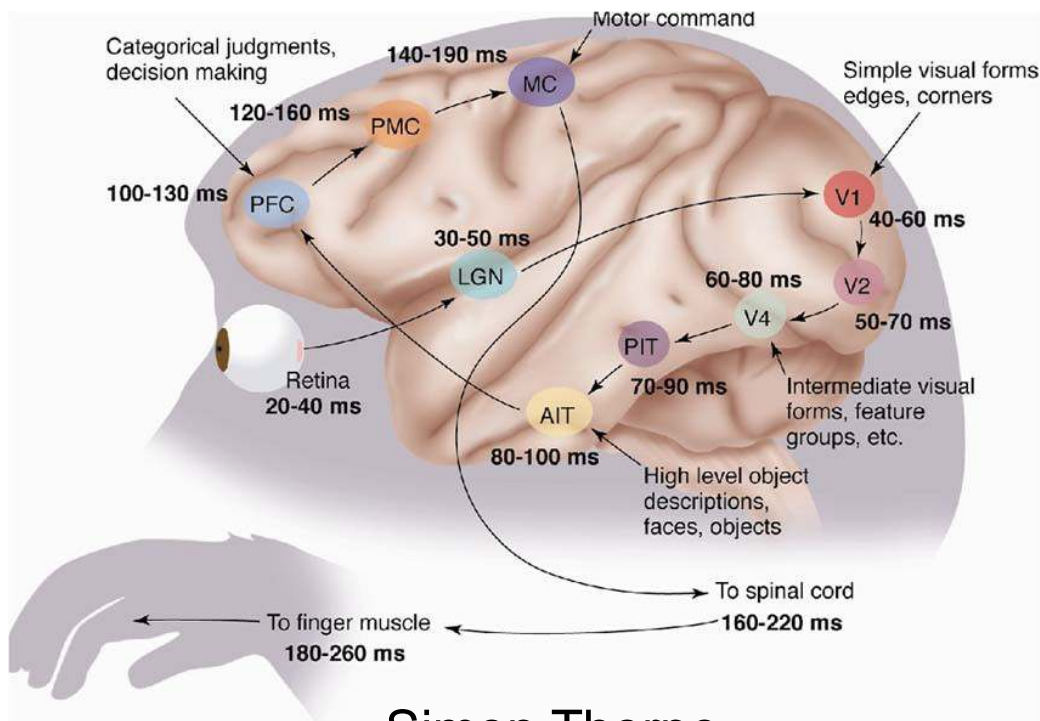
Filozófiai/evolúciós vonatkozás:

Az emberi agy marad legtovább plasztikus a születés után, s rendkívül hosszan megőrzi ezt a képességét a többi főemlőssel összehasonlítva

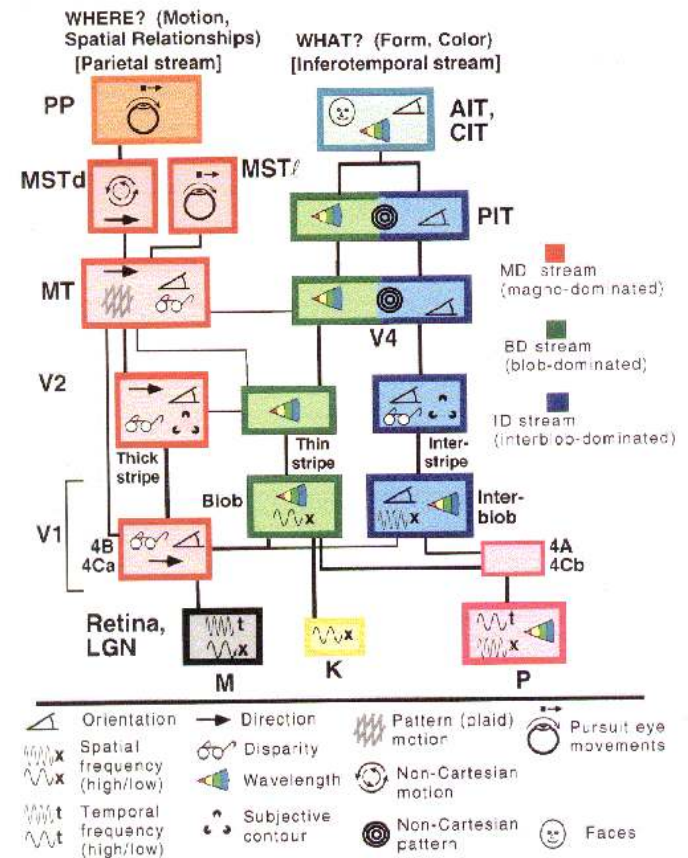
- ezáltal egy rugalmas, könnyen alkalmazkodó, újrachangolható rendszert kapunk
- Ez a rugalmasság védtelenséggel is jár, s akkori időszakban nagyobb védelemre van szükség

Hasonlít a biológiai struktúrákhoz

The visual cortex is hierarchical – the ventral (recognition) pathway has multiple stages

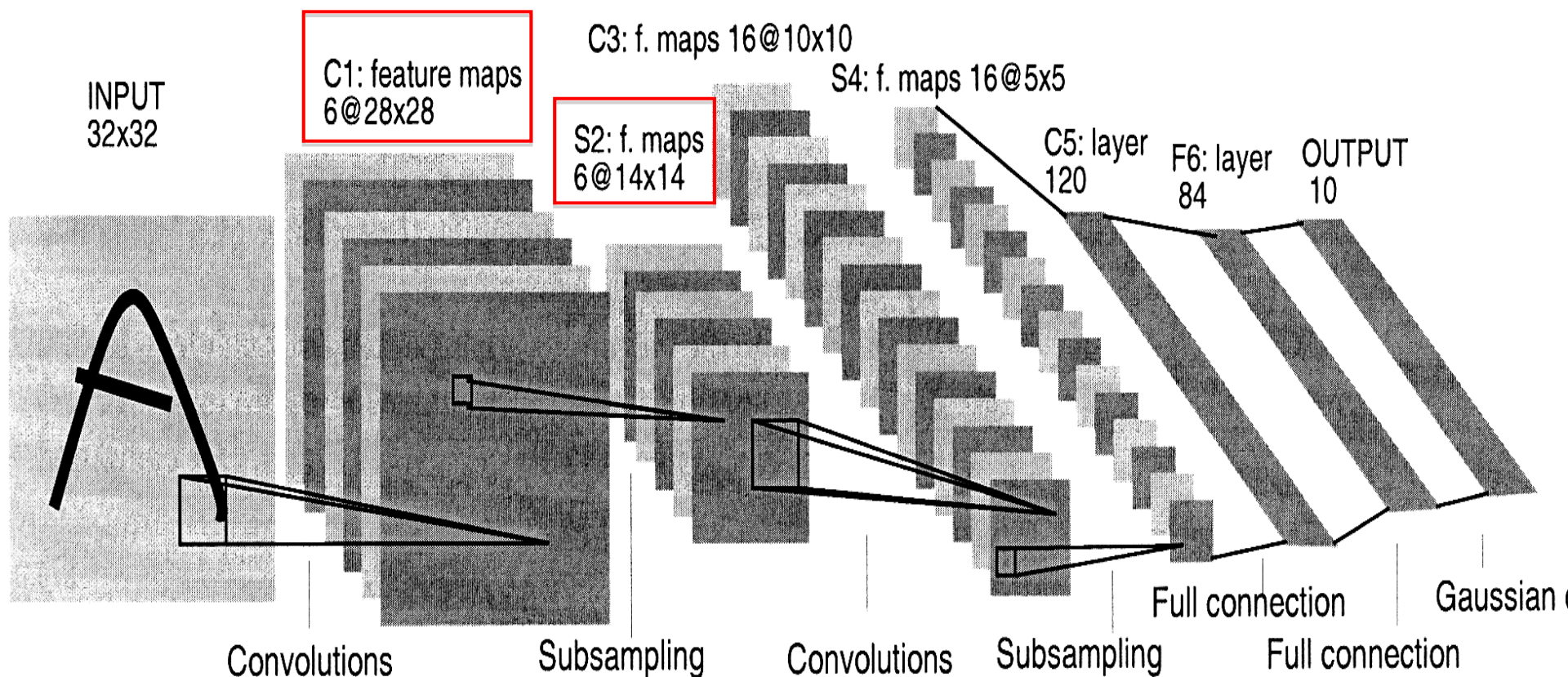


Simon Thorpe



Gallant & Van Essen

Egy gyakorlatban használt ConvNN struktúrája



A rendszer által hibásan osztályozott elemek

4	3	8	1	5	4	2	8	6	1
4→6	3→5	8→2	2→1	5→3	4→8	2→8	3→5	6→5	7→3

4	8	7	5	4	6	3	2	8	4
9→4	8→0	7→8	5→3	8→7	0→6	3→7	2→7	8→3	9→4

8	5	4	3	6	9	4	6	4	1
8→2	5→3	4→8	3→9	6→0	9→8	4→9	6→1	9→4	9→1

4	0	6	3	3	9	6	6	6	8
9→4	2→0	6→1	3→5	3→2	9→5	6→0	6→0	6→0	6→8

4	7	4	4	2	9	4	4	4	4
4→6	7→3	9→4	4→6	2→7	9→7	4→3	9→4	9→4	9→4

2	4	8	3	4	6	8	8	3	9
8→7	4→2	8→4	3→5	8→4	6→5	8→5	3→8	3→8	9→8

1	9	6	0	6	7	0	6	4	2
1→5	9→8	6→3	0→2	6→5	9→5	0→7	1→6	4→9	2→1

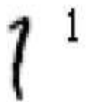
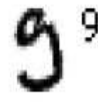
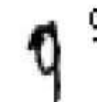
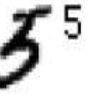
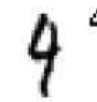
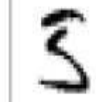
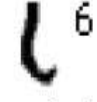
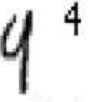
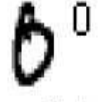
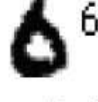
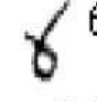
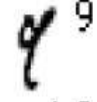
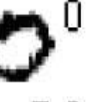
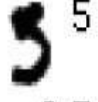
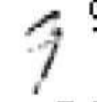
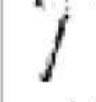
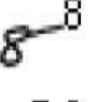
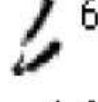
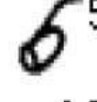
2	8	4	7	7	6	9	6	6	5
2→8	8→5	4→9	7→2	7→2	6→5	9→7	6→1	5→6	5→0

4	2
4→9	2→8

Összesen 82 hibás detekció

Az emberi hibák száma ezen a halmazon nagyjából 20-30 közötti

A rendszer által hibásan osztályozott elemek

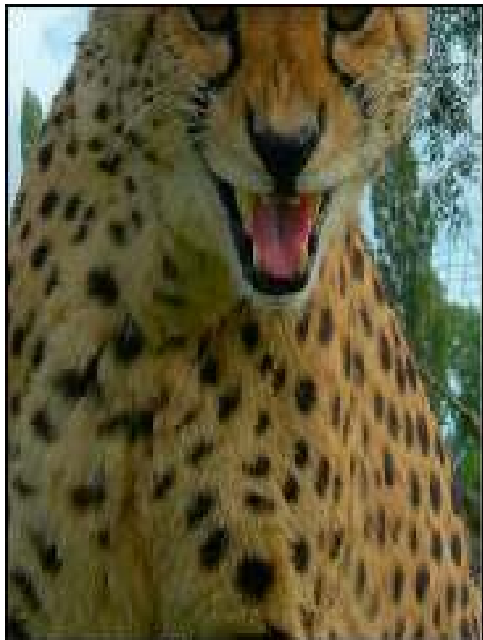
 17	 71	 98	 59	 79	 35	 23
 49	 35	 97	 49	 94	 02	 35
 16	 94	 60	 06	 86	 79	 71
 49	 50	 35	 98	 79	 17	 61
 27	 58	 78	 16	 65	 94	 60

A legfelső sorban található a legnagyobb súlyú válasz
Az alsó sorban a következő két válasz látható

A helyes válasz szinte mindig megtalálható az első három legvalószínűbb válaszban

Jelenleg 20-25 közti hibaszámot érnek el ezen az adathalmazon ~ megegyezi az emberi felismerés potosságával

Néhány példa a GoogLeNet esetében



cheetah

cheetah

leopard

snow leopard

Egyptian cat



bullet train is like a plane; with in-train magazine and a jacket that you can plug your hands/feet into and listen to

bullet train

bullet train

passenger car

subway train

electric locomotive



hand glass

scissors

hand glass

frying pan

stethoscope



mite

container ship

motor scooter

leopard

mite	container ship	motor scooter	leopard
black widow	lifeboat	go-kart	jaguar
cockroach	amphibian	moped	cheetah
tick	fireboat	bumper car	snow leopard
starfish	drilling platform	golfcart	Egyptian cat



grille

mushroom

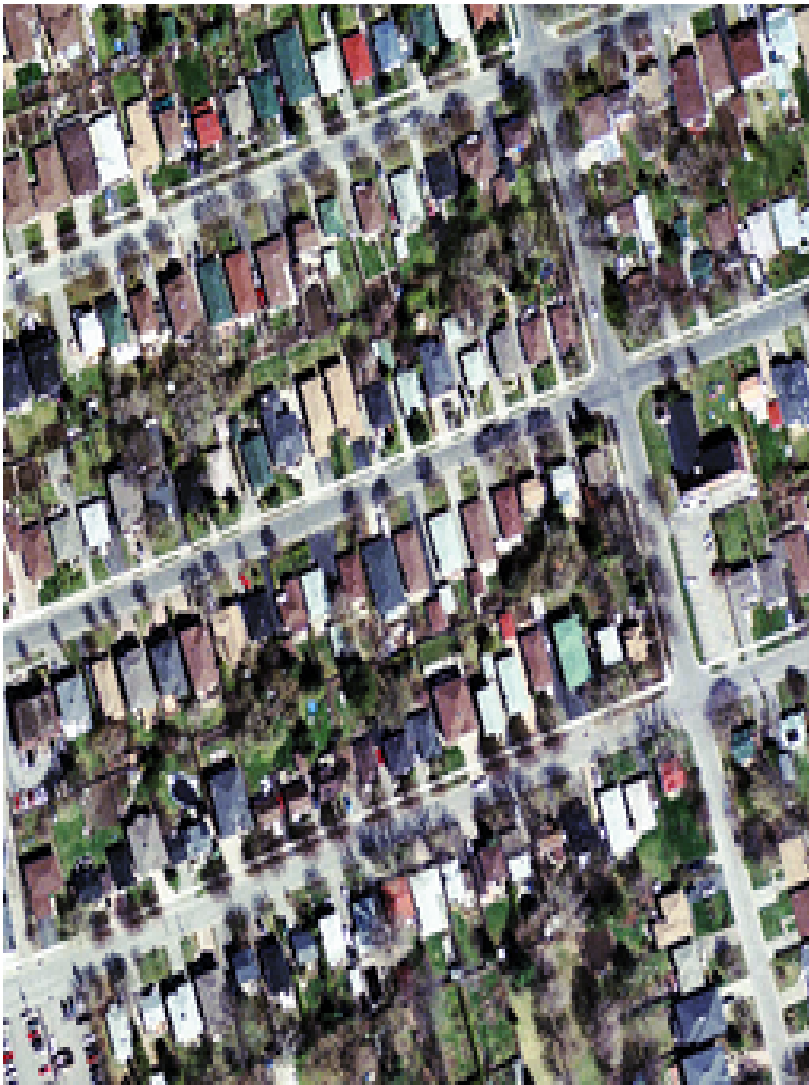
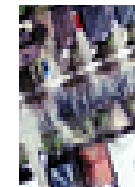
cherry

Madagascar cat

convertible	agaric	dalmatian	squirrel monkey
grille	mushroom	grape	spider monkey
pickup	jelly fungus	elderberry	titi
beach wagon	gill fungus	ffordshire bullterrier	indri
fire engine	dead-man's-fingers	currant	howler monkey

Sok esetben a hibákból látható, hogy a rendszer valóban általános tulajdonságokat tanul meg

Jelenleg a legjobb út azonosító algoritmus



A ConvNN hardware támogatás

A sokmagos architektúrákon, különösen GPU-k esetében ezek a számítások könnyen futtathatók (az architektúra konvolúciók számítására lett kitalálva SIMD)

Nvidia GTX 580 Graphics Processor Units (több, mint 1000 mag)

- Konvolúció, mint mátrix-mátrix szorzás hatékonyan számolható.
- GPU-s nagy sávszélesség a memóriához
- A tanítás sokszor meglehetősen lassú → 1 hét is lehet.

A felismerés 'néhány' konvolúcióval kiszámolható → gyakorlatilag valós időben

Sparsity - ritkaság

Egy Sparse reprezentáció segít a tanulásban. Ha már a közbűnső rétegekben figyelembe vesszük, hogy 'saprse' reprezentációt szeretnénk – azaz csak néhány feature esetében akarunk nagy választ, akkor jelentősen javíthatjuk a detekciót.

Az itt látható lenyomatok közül ezzel specifikusabbakat választhatunk ki, s érezhető, hogy jobban használható egy olyan minta, ahol 1-2 erős választ kapunk, s a többi csatorna válasza szinte mind nulla, mint egy olyan, ahol minden csatornán találunk választ.

Ez a 'sparsity' az emberi retina csatornáiban is megtalálható

