

COMP541

DEEP LEARNING

Lecture #01 – Introduction



**KOÇ
UNIVERSITY**

Aykut Erdem // Koç University // Fall 2025

Welcome to COMP541

- This course gives an overview of deep learning,
- In particular, we will cover various deep architectures and deep learning methods.
- You will develop fundamental and practical skills at applying deep learning to your research.

Updated in Fall
2024 to cover LLMs
in more depth!

A little about me...

Koç University
Associate Professor
2020-now



Hacettepe University
Associate Professor
2010-2020



Università Ca' Foscari di Venezia
Post-doctoral Researcher
2008-2010



Middle East Technical University
1997-2008
Ph.D., 2008
M.Sc., 2003
B.Sc., 2001



MIT
Fall 2007
Visiting Student



VirginiaTech
Visiting Research Scholar
Summer 2006



The broad goal of my research is to explore better ways to **understand**, **interpret**, and **manipulate** visual data.

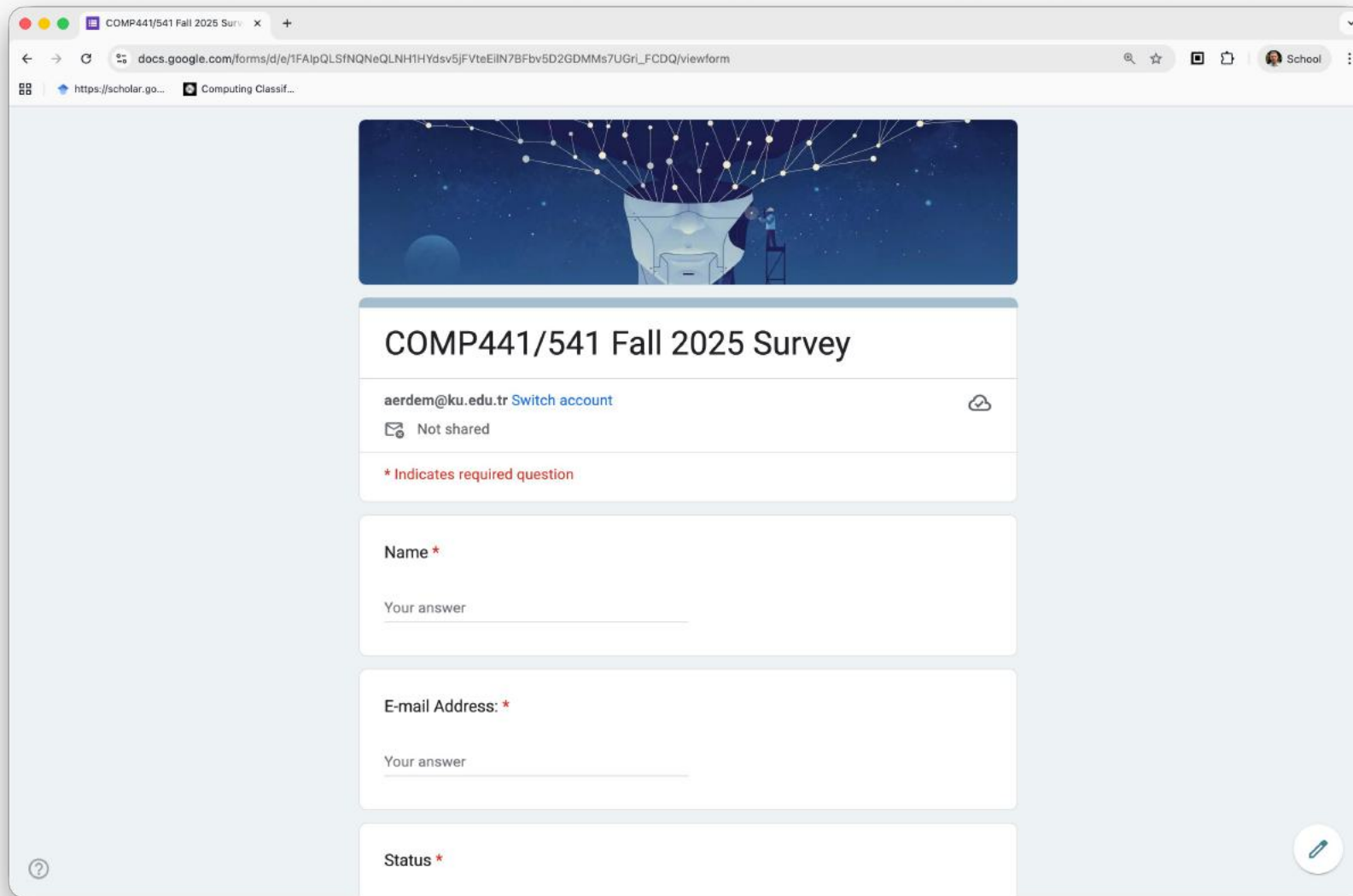
Research Interests

- Deep Learning
- Generative AI
- Computer Vision
- Language Understanding



 <https://aykuterdem.github.io>

Now, what about you?



The screenshot shows a Google Forms interface in a web browser. The browser's address bar displays the URL: `docs.google.com/forms/d/e/1FAIpQLSfNQNeQLNH1HYdsv5JFVteEiIN7BFbv5D2GDMMs7UGrI_FCDQ/viewform`. The form title is "COMP441/541 Fall 2025 Survey". Below the title, the user is identified as "aerdem@ku.edu.tr" with a "Switch account" link. A status indicator shows "Not shared". A red asterisk note states: "* Indicates required question". The form contains three visible questions: "Name *" with a text input field labeled "Your answer"; "E-mail Address: *" with a text input field labeled "Your answer"; and "Status *" with a text input field. A blue pencil icon for editing is located in the bottom right corner of the form area.

<https://forms.gle/sHgeXiYiPQaqkcGy7>



Course Logistics

Course Information

Lectures Tuesday and Thursday 14:30-15:40 (SNA 159)

PS Thursday 17:30-18:40 (CASE Z08)

Instructor Aykut Erdem

TAs Andrew Bond & Hakan Capuk.

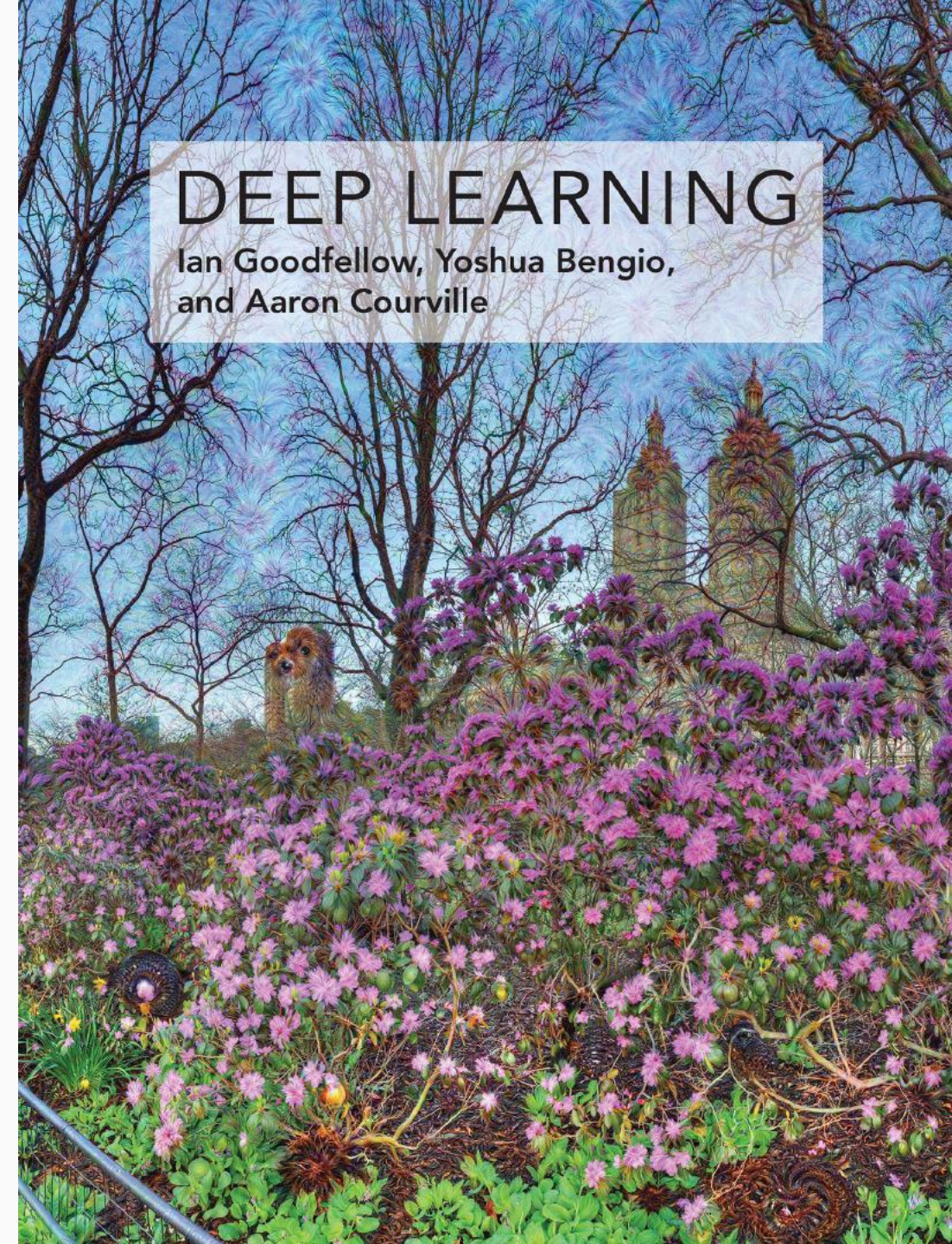


Website <https://aykuterdem.github.io/classes/comp541.f25/>

- KUHub Learn for course related announcements and collecting and grading your submissions

Textbook

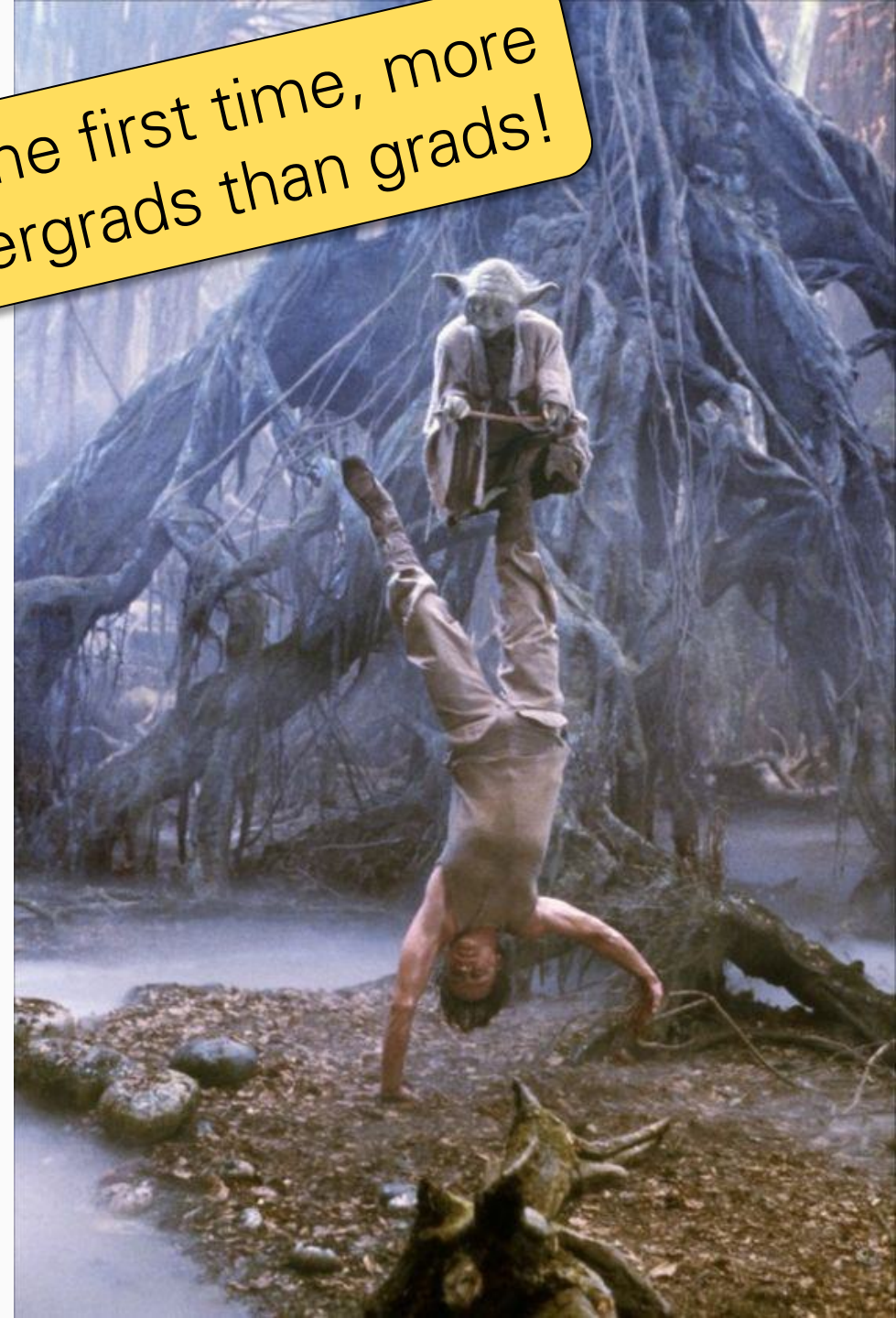
- Goodfellow, Bengio, and Courville, Deep Learning, MIT Press, 2016 (draft available [online](#))
- In addition, we will extensively use online materials (video lectures, blog posts, surveys, papers, etc.)



Instruction style

- Students are responsible for studying and keeping up with the course material outside of class time.
 - Reading particular book chapters, papers or blogs, or
 - Watching some video lectures.
- After the first four lectures, each week students will present papers related to the topics of the previous week.
 - Weekly paper reviews will be prepared by all the students

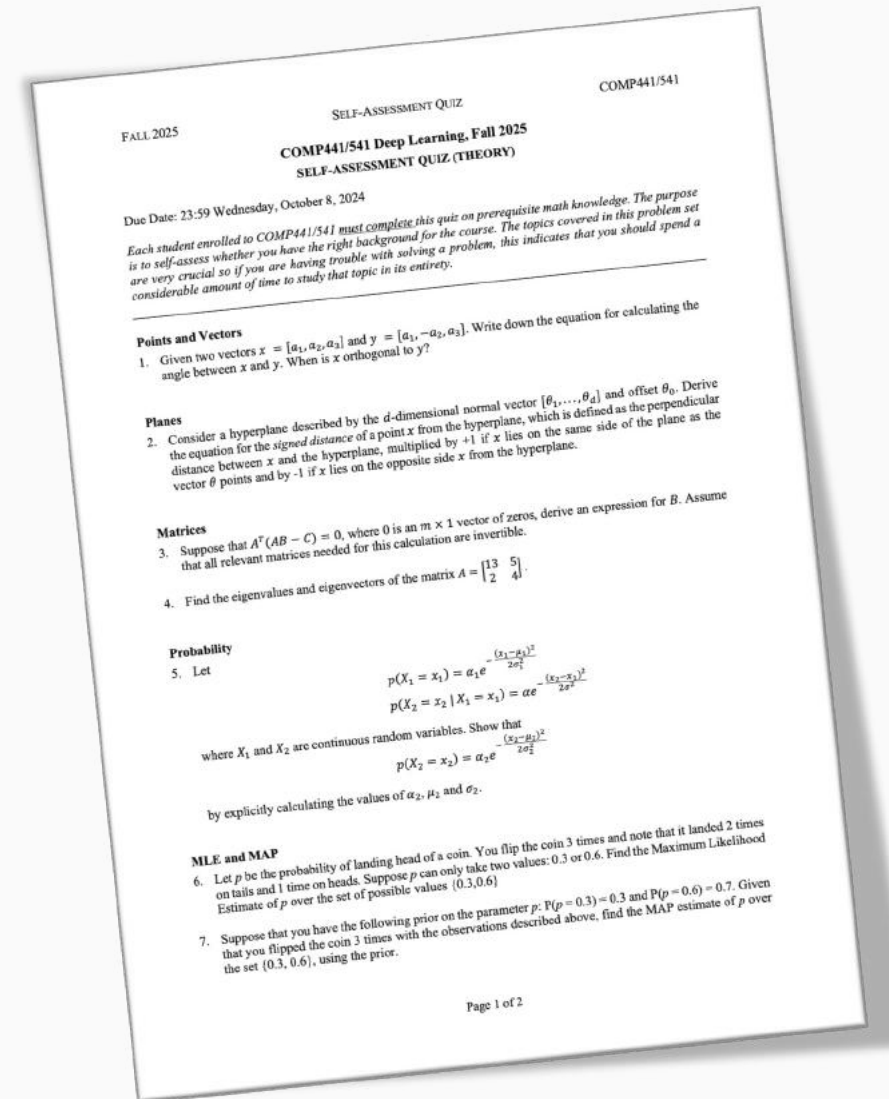
For the first time, more undergrads than grads!



Prerequisites

- Calculus and linear algebra
 - Derivatives,
 - Matrix operations
- Probability and statistics
- Machine learning
- Programming

Read Chapter 2-4
of the Deep Learning textbook for a quick review.



Self-Assessment Quiz (Theory)

Due Date: October 8 (23:59).

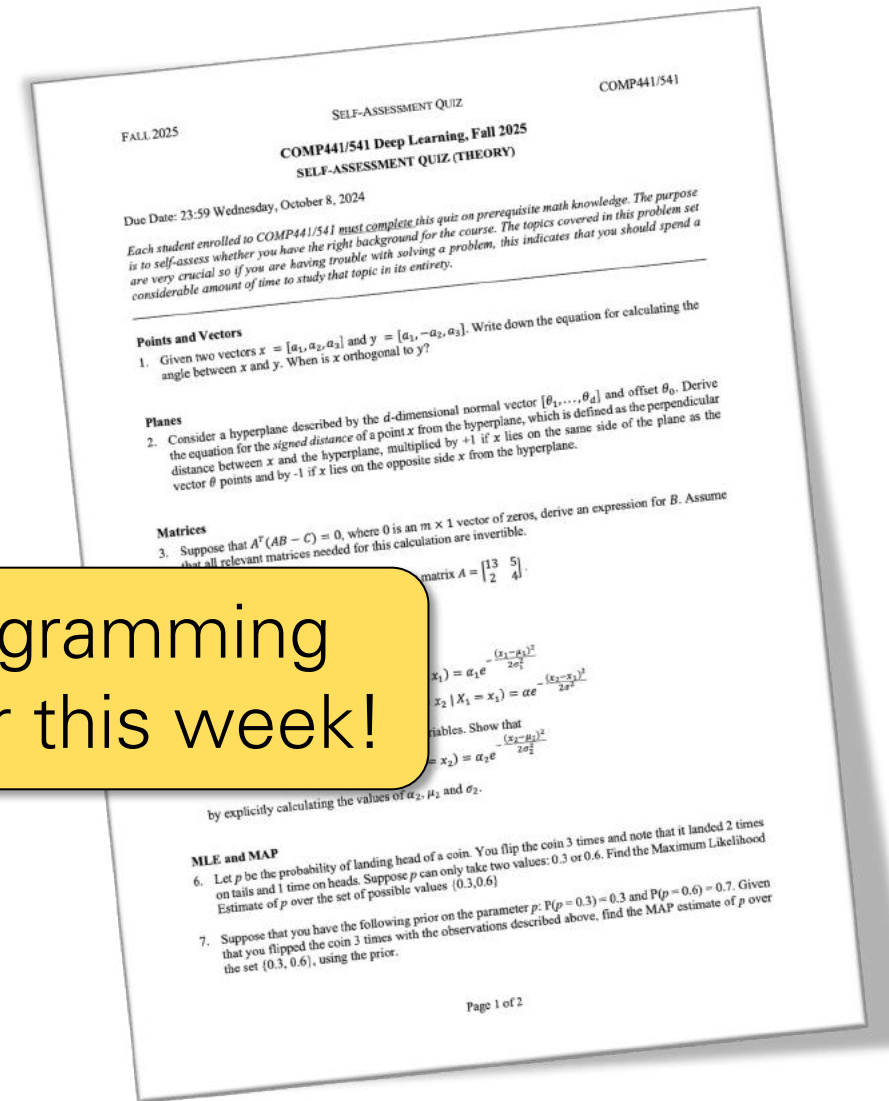
Each student enrolled to COMP441/541 must complete and pass this quiz!

Prerequisites

- Calculus and linear algebra
 - Derivatives,
 - Matrix operations
- Probability
- Machine learning
- Programming

The self-assessment quiz on programming background will be released later this week!

Read Chapter 2-4
of the Deep Learning textbook for a quick review.



Self-Assessment Quiz (Theory)

Due Date: October 8 (23:59).

Each student enrolled to COMP441/541 must complete and pass this quiz!

Topics Covered in ENGR 421

- **Basics of Statistical Learning**

- Loss function, MLE, MAP, Bayesian estimation, bias-variance tradeoff, overfitting, regularization, cross-validation

- **Supervised Learning**

- Nearest Neighbor, Naïve Bayes, Logistic Regression, Support Vector Machines, Kernels, Neural Networks, Decision Trees
- Ensemble Methods: Bagging, Boosting, Random Forests

- **Unsupervised Learning**

- Clustering: K-Means, Gaussian mixture models
- Dimensionality reduction: PCA, SVD

Grading

Self-Assessment Quiz	2%
Programming Assignments	20% (4 assignments x 5% each)
Midterm Exam	17%
Course Project	36%
Paper Presentations	10%
Paper Reviews	5%
Class Participation	10%

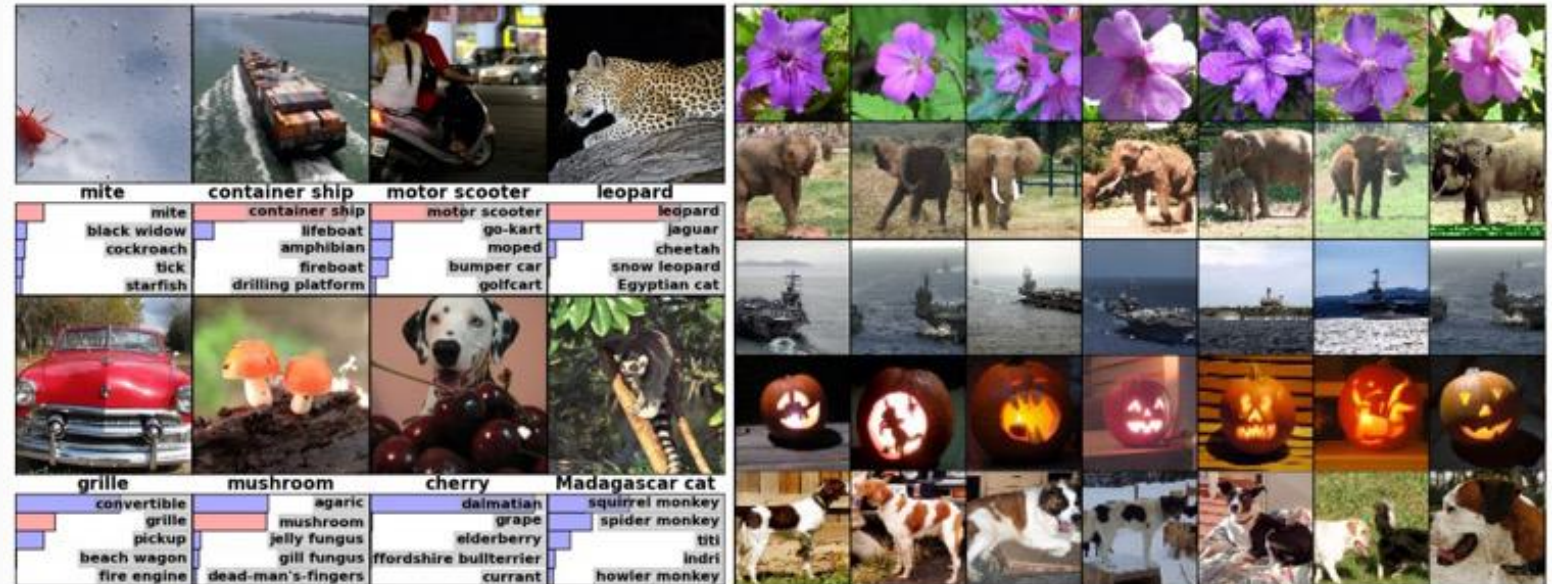
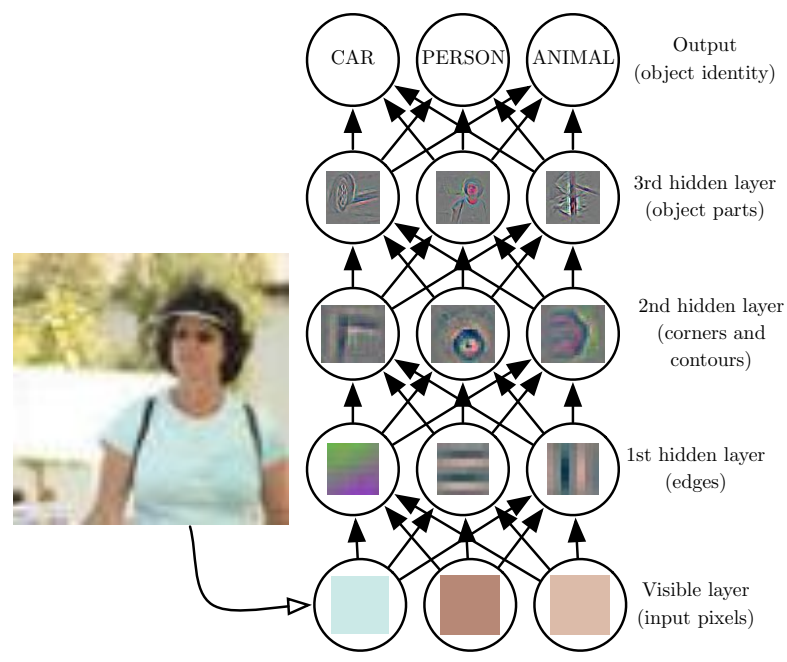
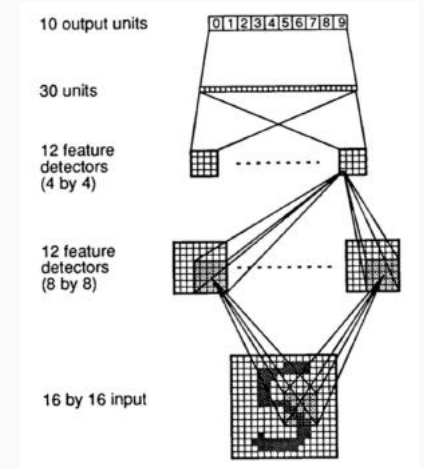
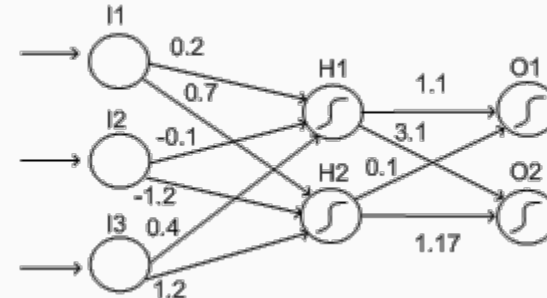
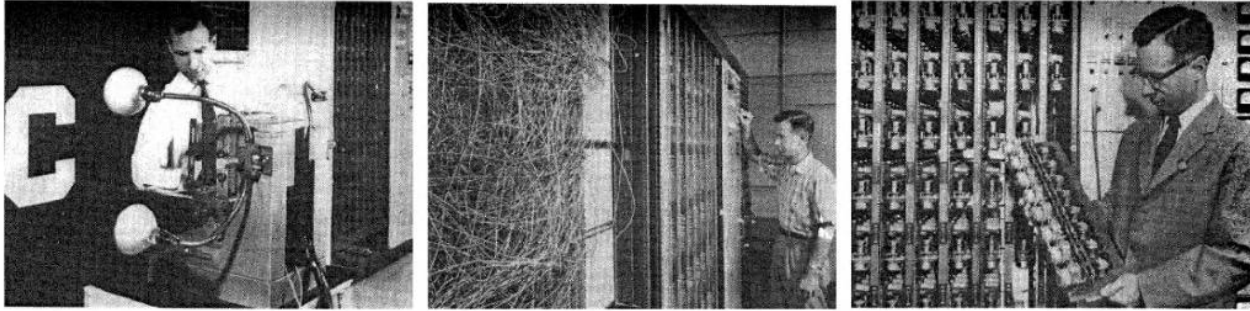
Schedule

Week 1	Introduction to Deep Learning
Week 2	Machine Learning Overview
Week 3	Multi-Layer Perceptrons
Week 4	Training Deep Neural Networks
Week 5	Convolutional Neural Networks
Week 6	Understanding and Visualizing CNNs
Week 7	Recurrent Neural Networks

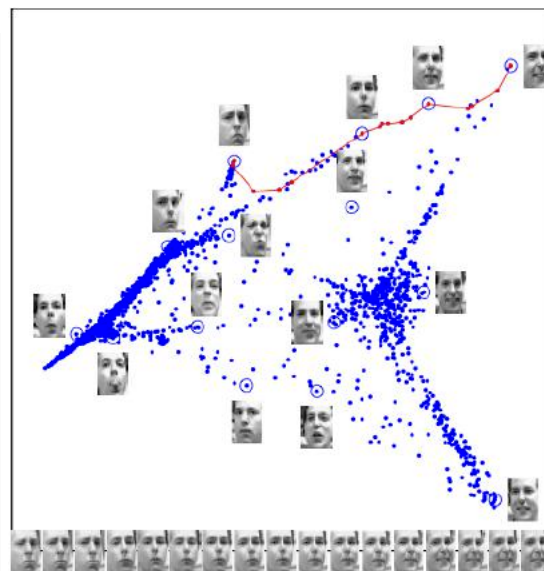
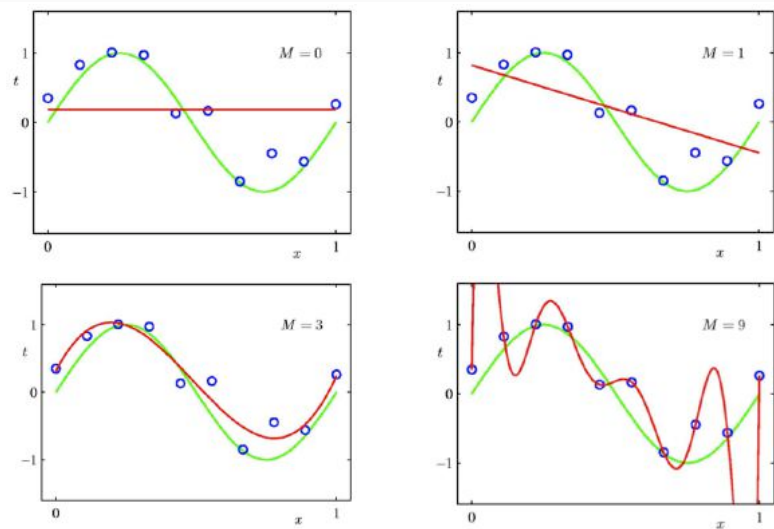
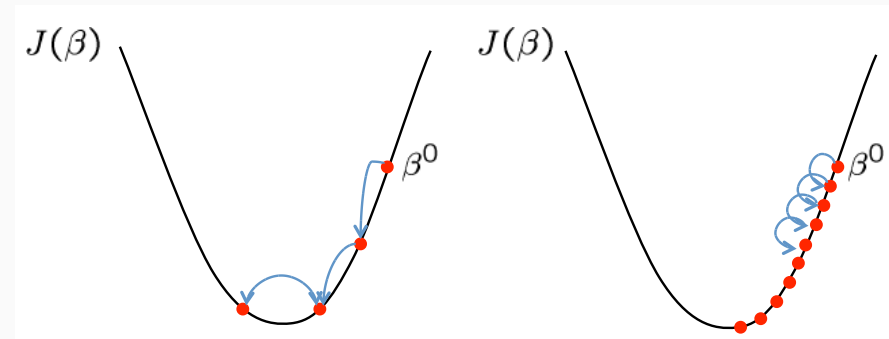
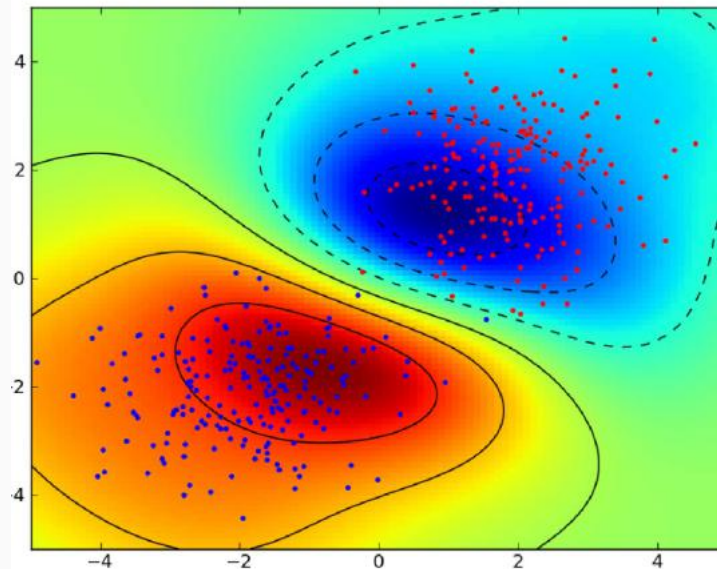
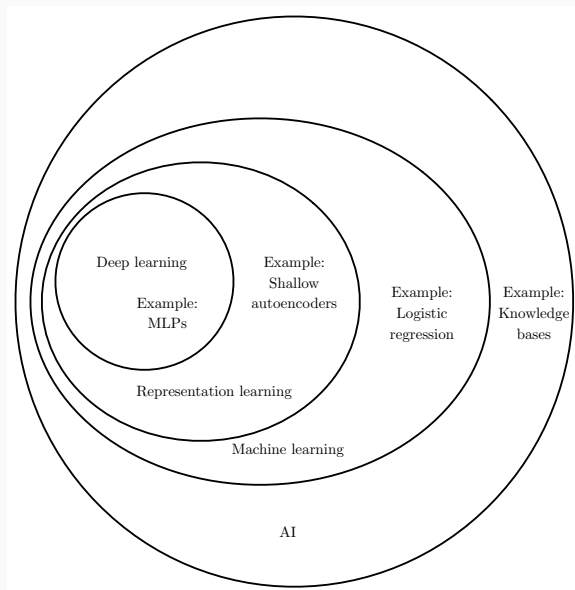
Schedule

Week 8	Attention and Transformers
Week 9	Graph Neural Networks
Week 10	Language Model Pretraining
Week 11	Project Progress Presentations
Week 12	Large Language Models
Week 13	Adapting LLMs
Week 14	Multimodal Pretraining

Lecture 1: Introduction to Deep Learning

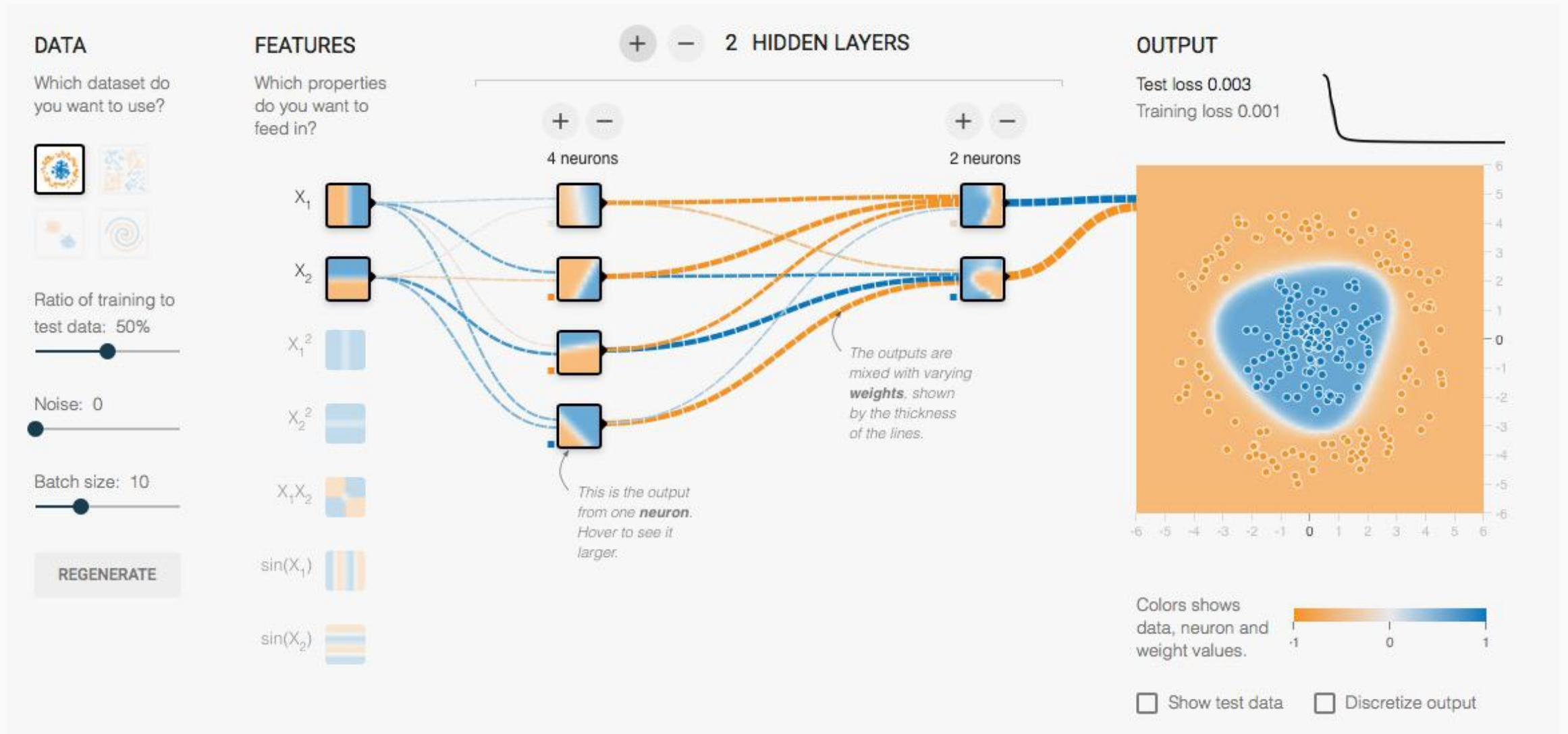


Lecture 2: Machine Learning Overview

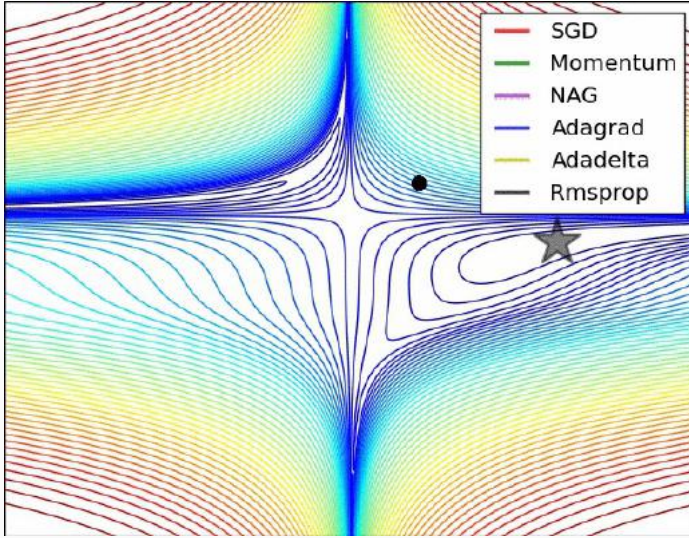


8	9	0	1	2	3	4	7	8	9	0	1	2	3	4	5	6	7	8	6
4	2	6	4	7	5	5	4	7	8	9	2	9	3	9	3	8	2	0	5
0	1	0	4	2	6	5	3	5	3	8	0	0	3	4	1	5	3	0	8
3	0	6	2	7	1	1	8	1	7	1	3	8	9	7	6	7	4	1	6
7	5	1	7	1	9	8	0	6	9	4	9	9	3	7	1	9	2	2	5
3	7	8	2	3	4	5	6	7	8	9	0	1	2	3	4	5	6	7	0
1	2	3	4	5	6	7	8	9	8	1	0	5	5	1	9	0	4	1	9
3	8	4	7	7	8	5	0	6	5	5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

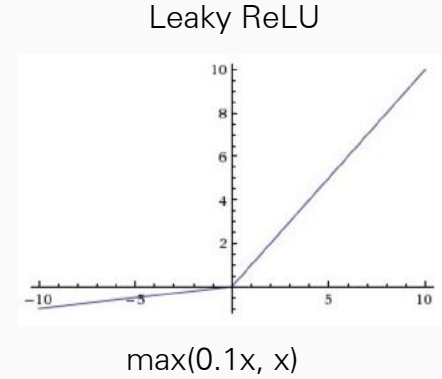
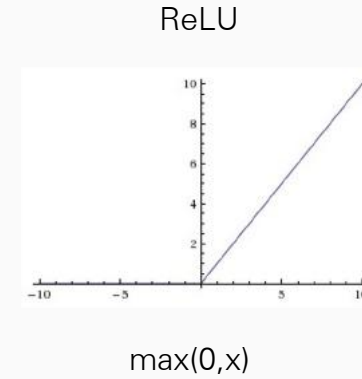
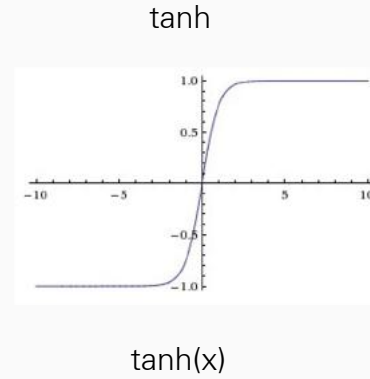
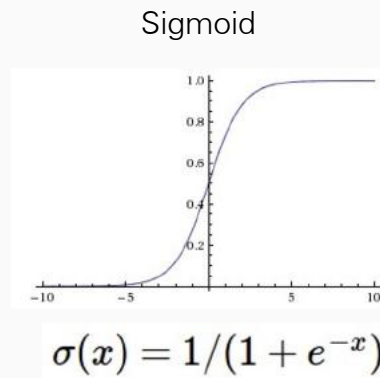
Lecture 3: Multi-Layer Perceptrons



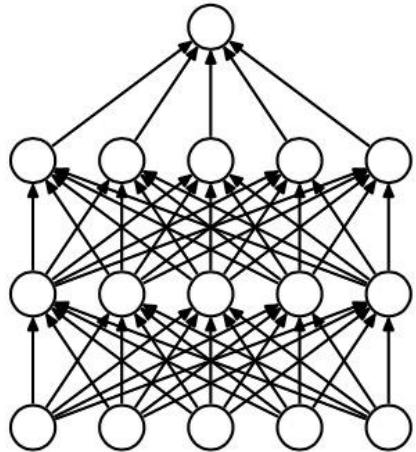
Lecture 4: Training Deep Neural Networks



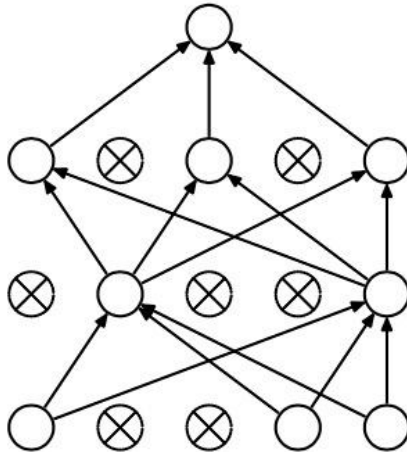
Optimizers



Activation Functions



(a) Standard Neural Net



(b) After applying dropout.

Dropout

Input: Values of x over a mini-batch: $\mathcal{B} = \{x_1 \dots x_m\}$;

Parameters to be learned: γ, β

Output: $\{y_i = \text{BN}_{\gamma, \beta}(x_i)\}$

$$\mu_{\mathcal{B}} \leftarrow \frac{1}{m} \sum_{i=1}^m x_i \quad // \text{ mini-batch mean}$$

$$\sigma_{\mathcal{B}}^2 \leftarrow \frac{1}{m} \sum_{i=1}^m (x_i - \mu_{\mathcal{B}})^2 \quad // \text{ mini-batch variance}$$

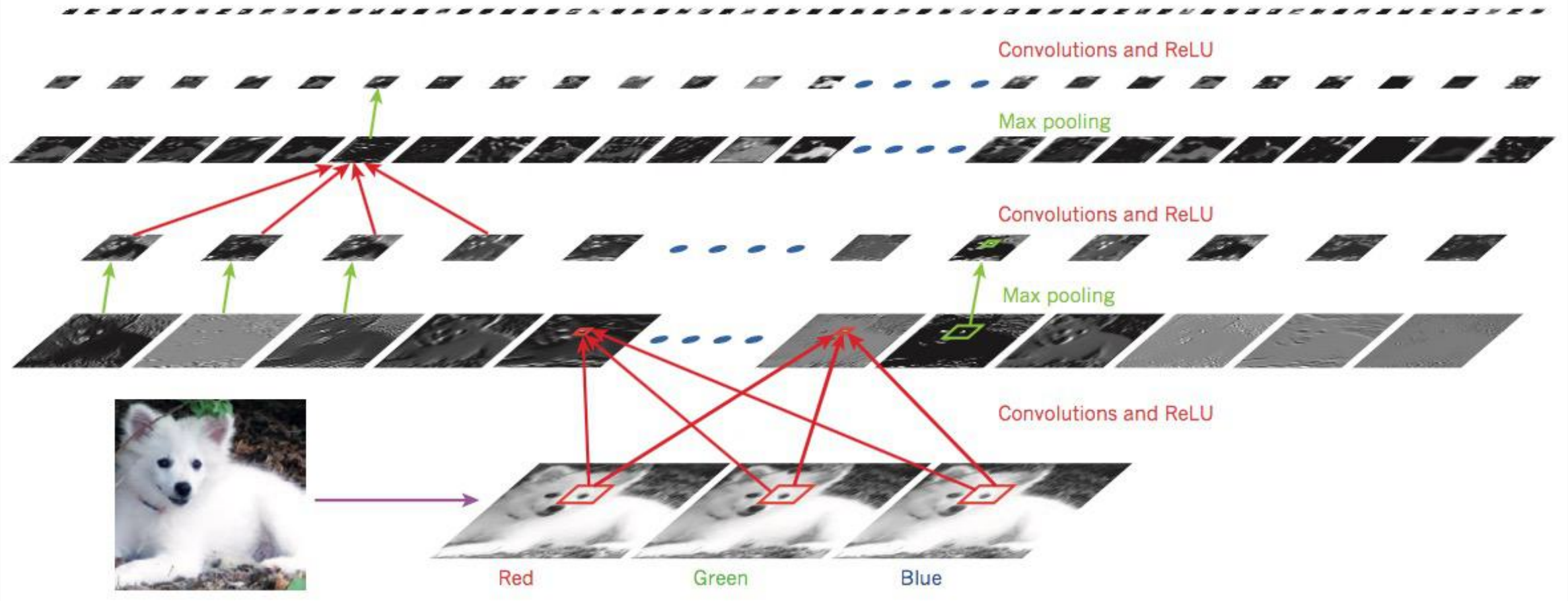
$$\hat{x}_i \leftarrow \frac{x_i - \mu_{\mathcal{B}}}{\sqrt{\sigma_{\mathcal{B}}^2 + \epsilon}} \quad // \text{ normalize}$$

$$y_i \leftarrow \gamma \hat{x}_i + \beta \equiv \text{BN}_{\gamma, \beta}(x_i) \quad // \text{ scale and shift}$$

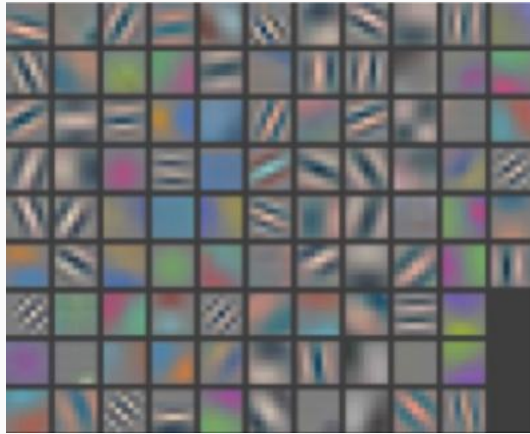
Batch Normalization

Lecture 5: Convolutional Neural Networks

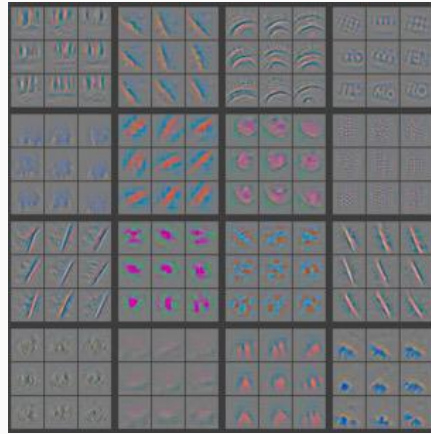
Samoyed (16); Papillon (5.7); Pomeranian (2.7); Arctic fox (1.0); Eskimo dog (0.6); white wolf (0.4); Siberian husky (0.4)



Lecture 6: Understanding and Visualizing CNNs



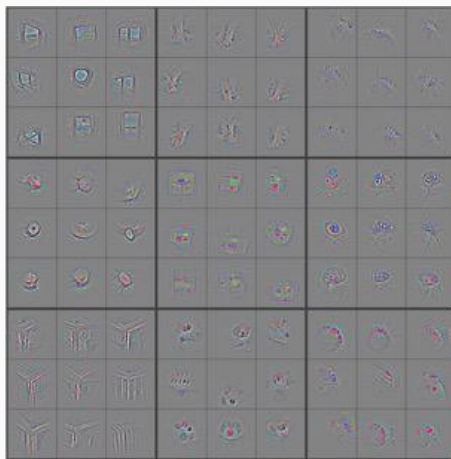
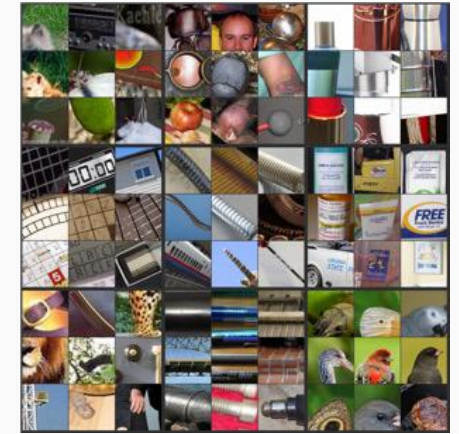
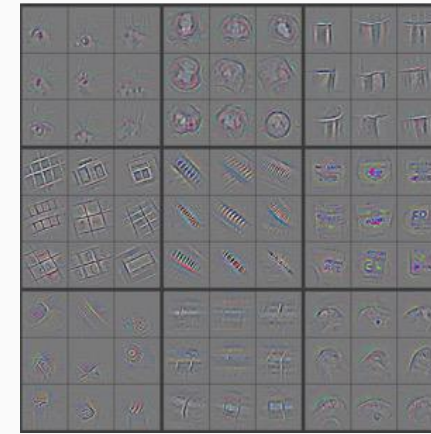
Layer 1



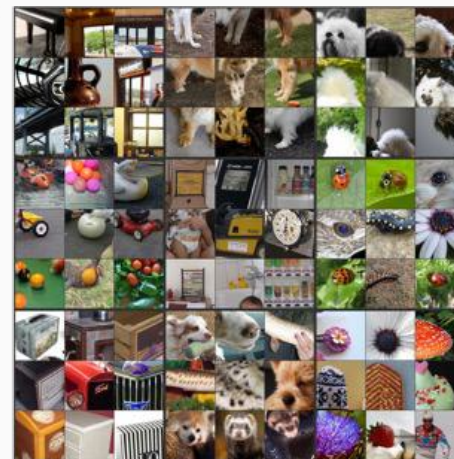
Layer 2



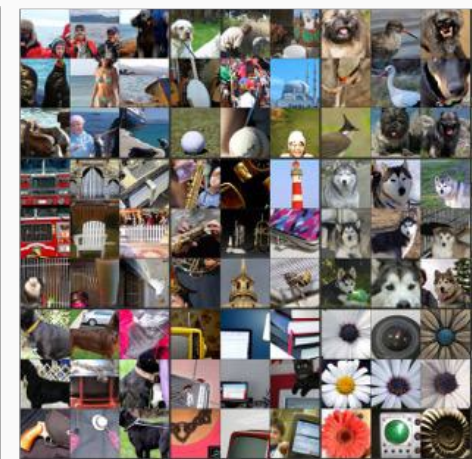
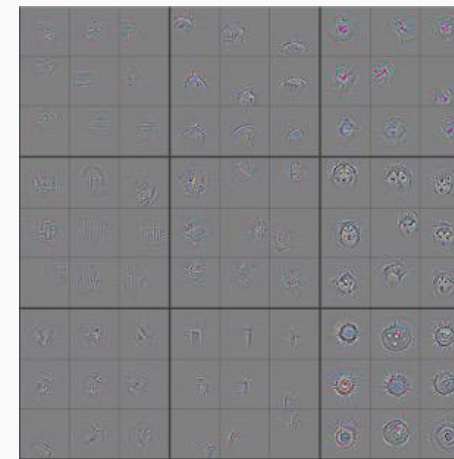
Layer 3



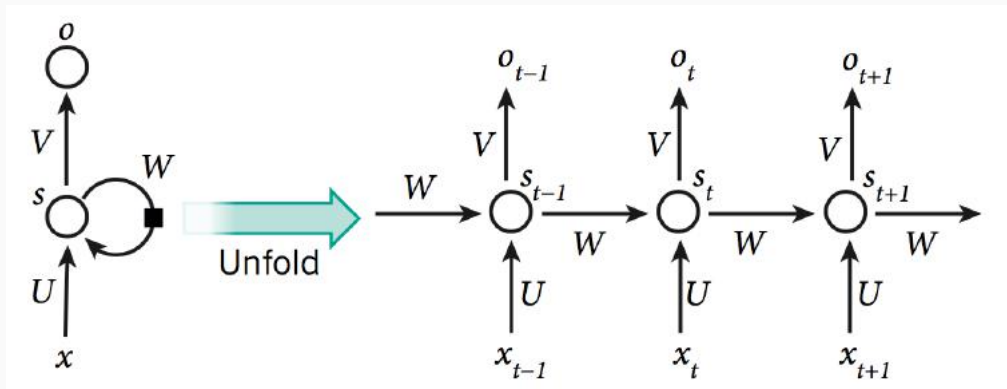
Layer 4



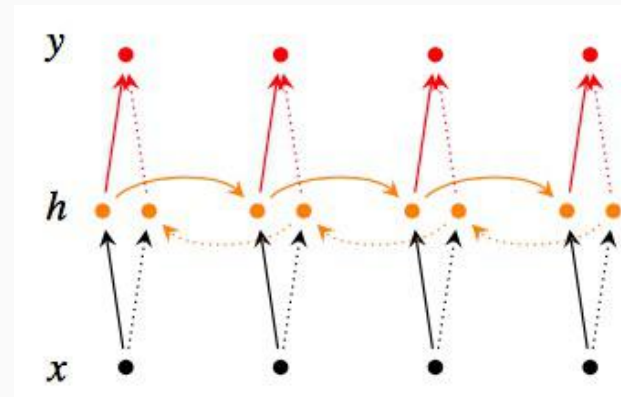
Layer 5



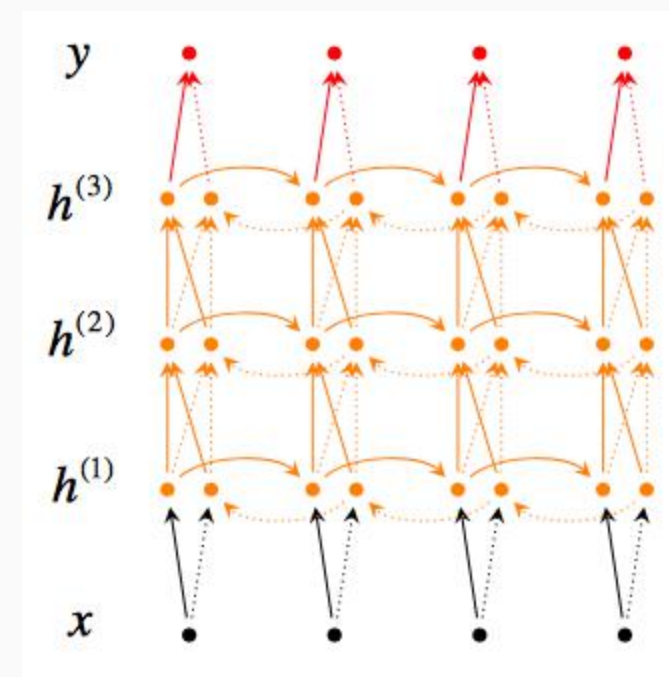
Lecture 7: Recurrent Neural Networks



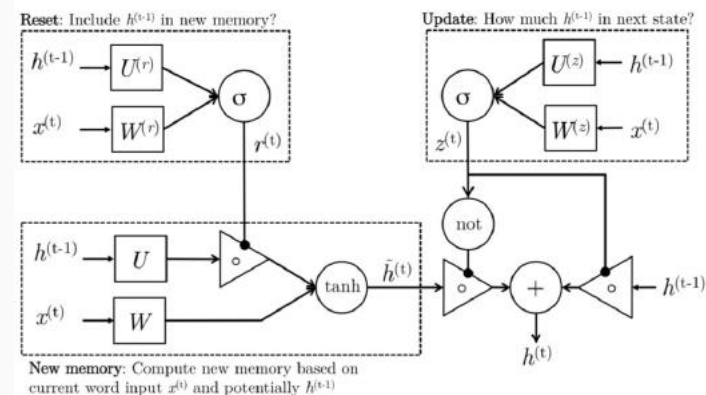
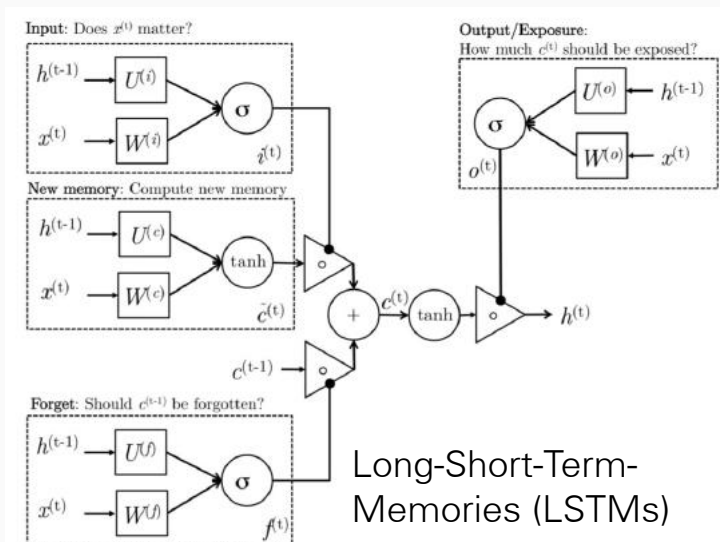
A Recurrent Neural Network (RNN)
(unfolded across time-steps)



A bi-directional RNN



A deep bi-directional RNN



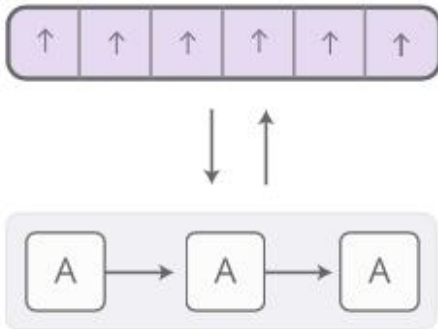
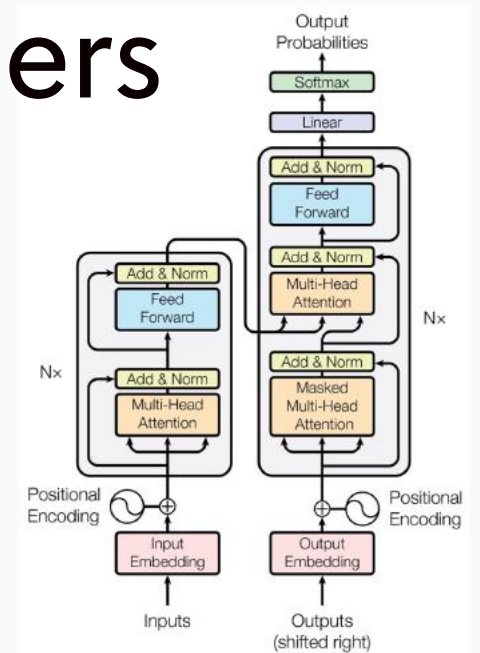
Lecture 8: Attention and Transformers



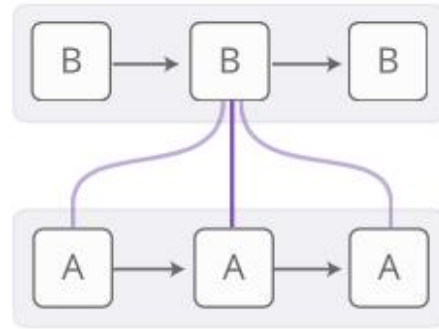
A little girl sitting on a bed with a teddy bear.



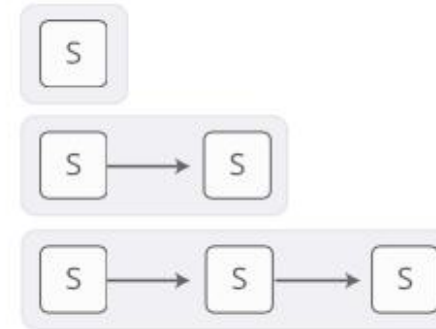
A group of people sitting on a boat in the water.



Neural Turing Machines

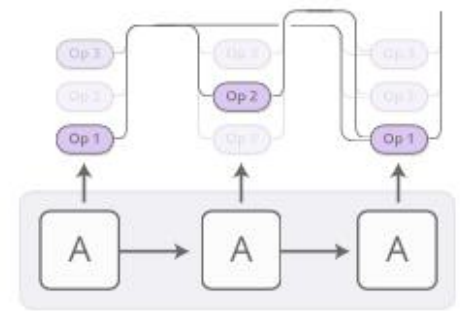


Attentional Interfaces



Adaptive Computation Time

Transformer Architecture



Neural Programmers

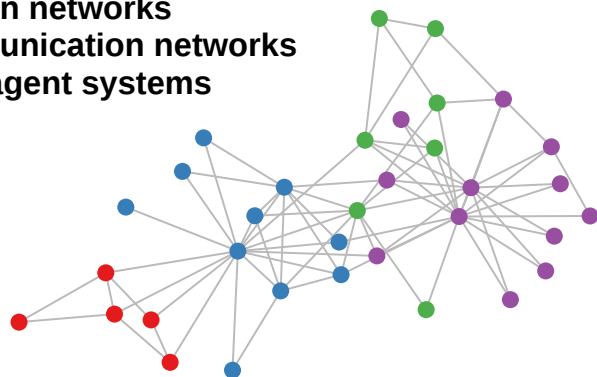
K. Xu et al., "Show, Attend and Tell: Neural Image Caption Generation with Visual Attention", ICML 2015

C. Olah and S. Carter, "Attention and Augmented Recurrent Neural Networks", Distill, 2016

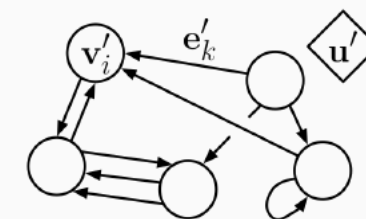
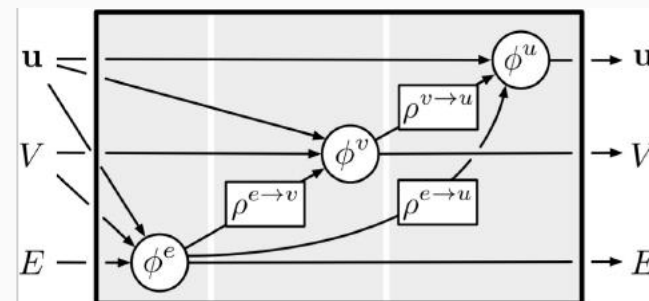
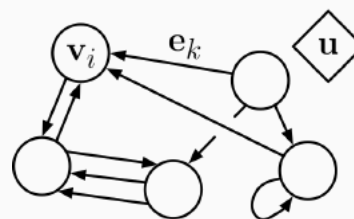
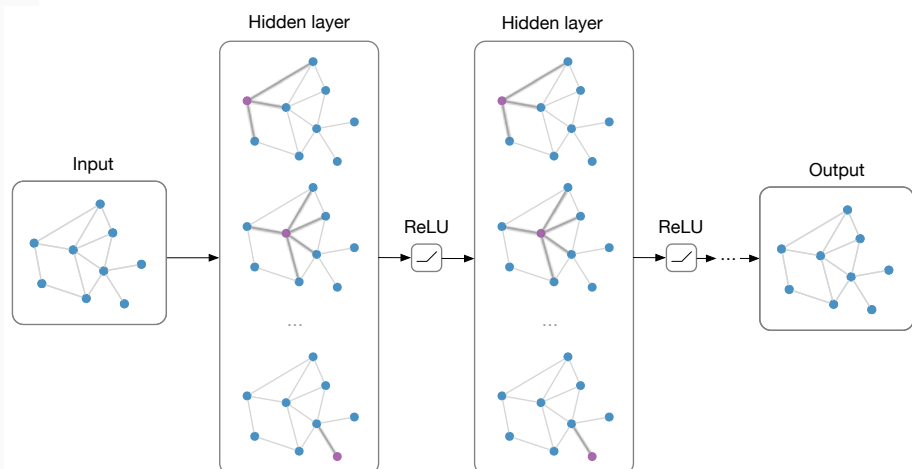
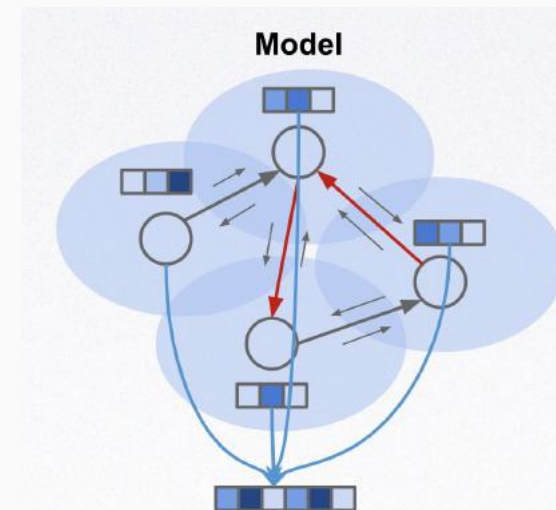
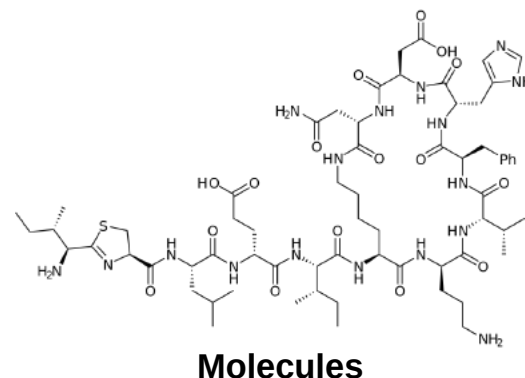
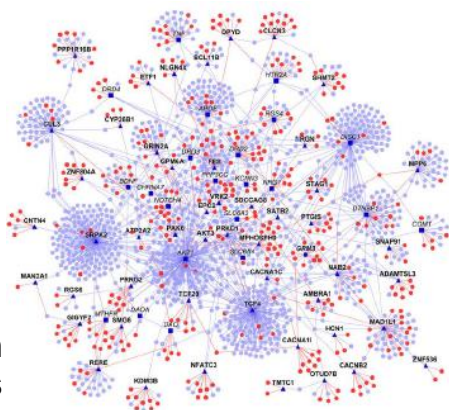
A. Vaswani et al. "Attention is All You Need", NeurIPS 2017.

Lecture 9: Graph Networks

Social networks
Citation networks
Communication networks
Multi-agent systems



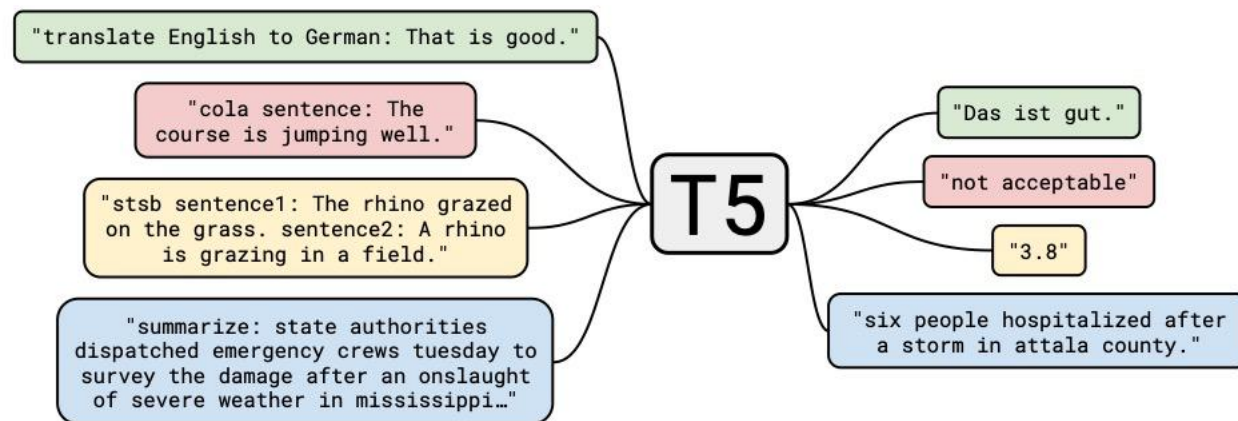
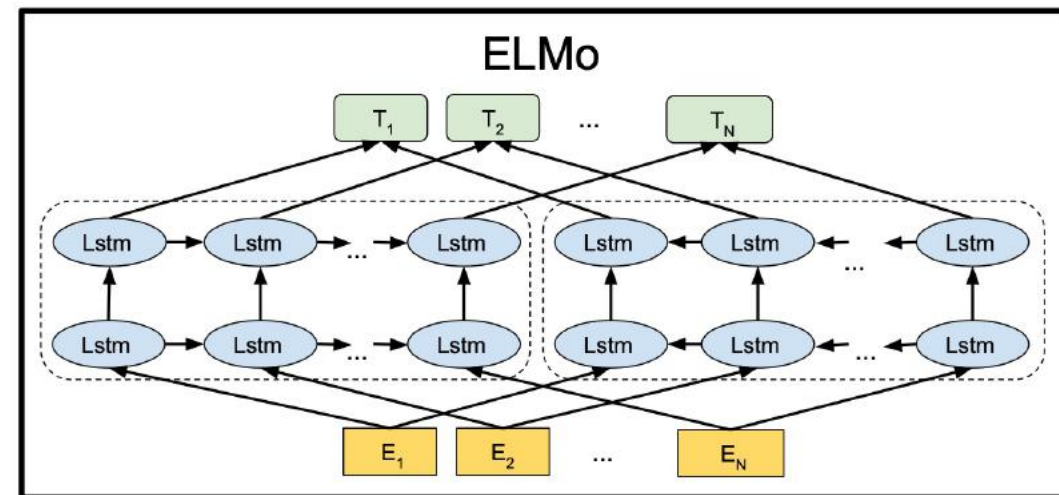
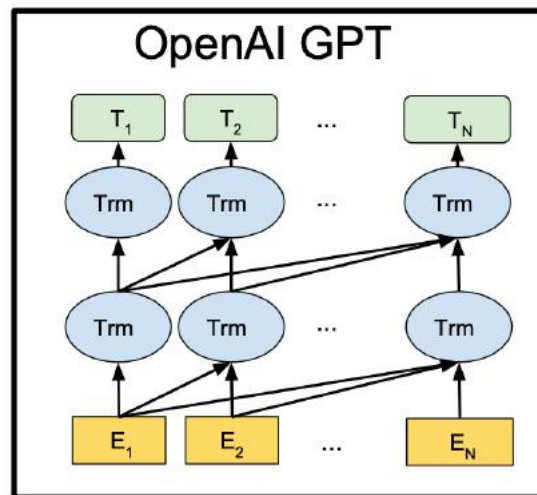
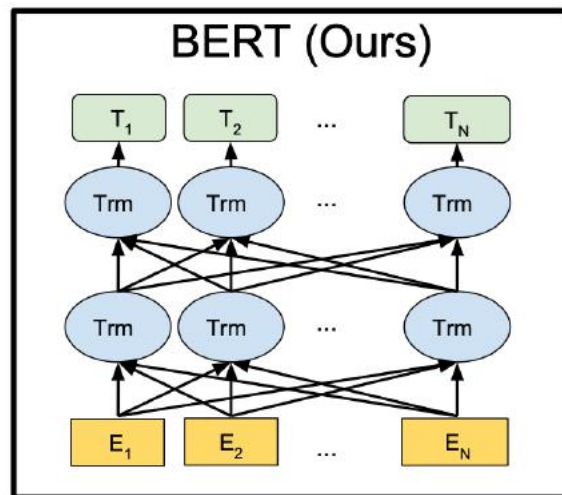
Protein interaction networks



T.N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks", ICLR 2017

P. Battaglia et al., "Relational inductive biases, deep learning, and graph networks", arXiv 2018

Week 10: Pretraining Language Models



Lecture 11: Large Language Models

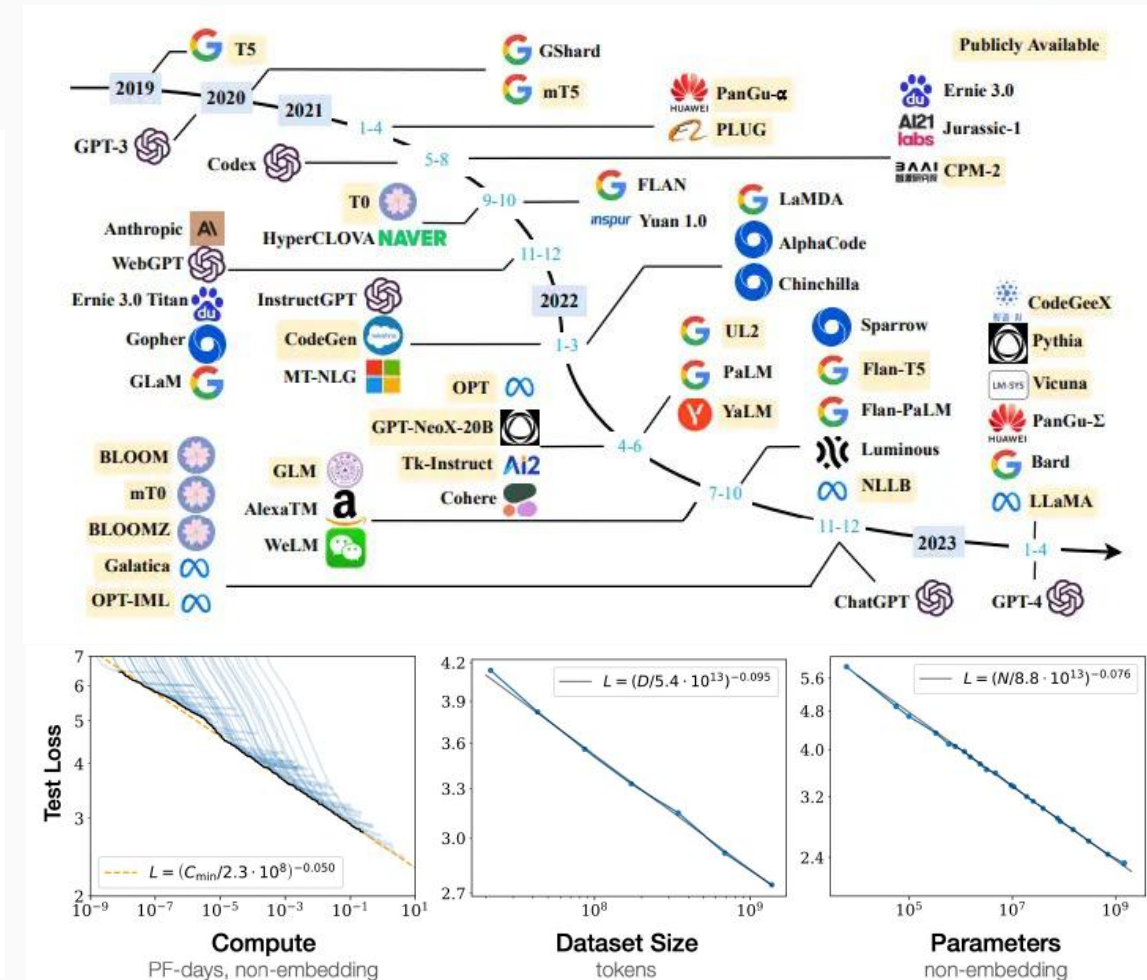
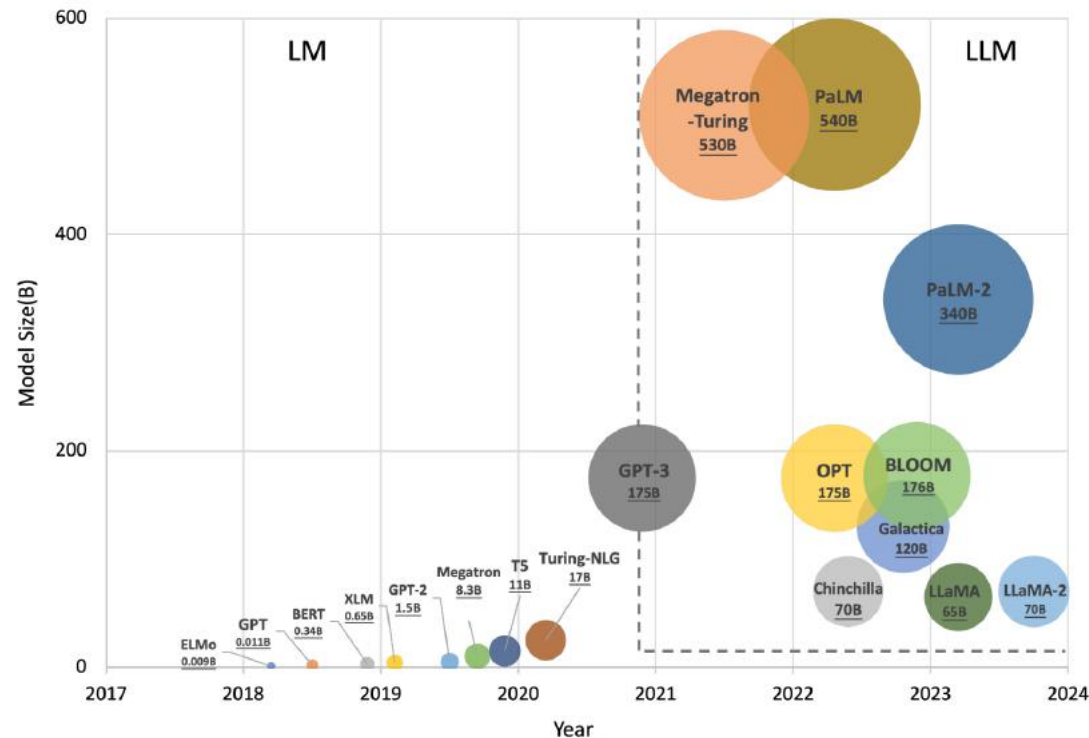
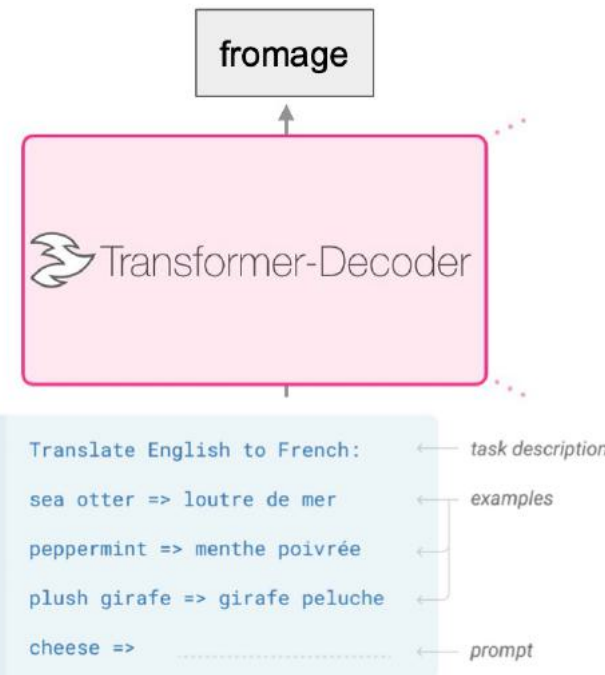
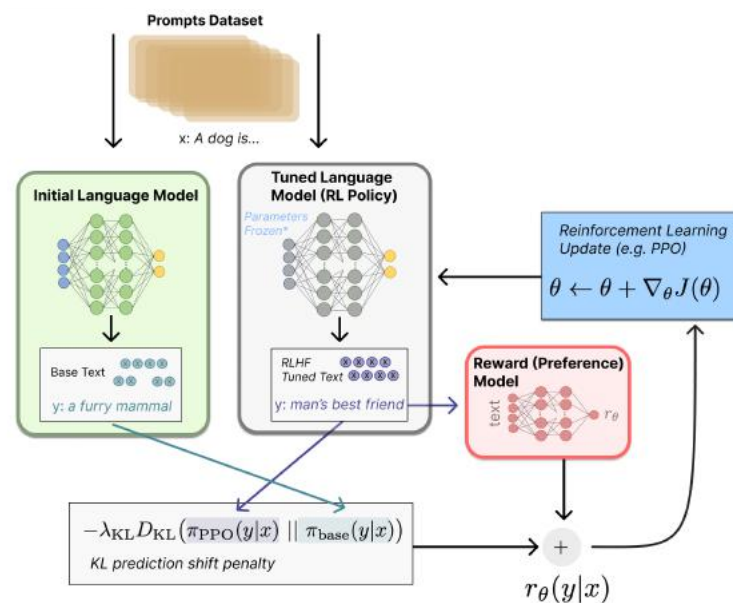
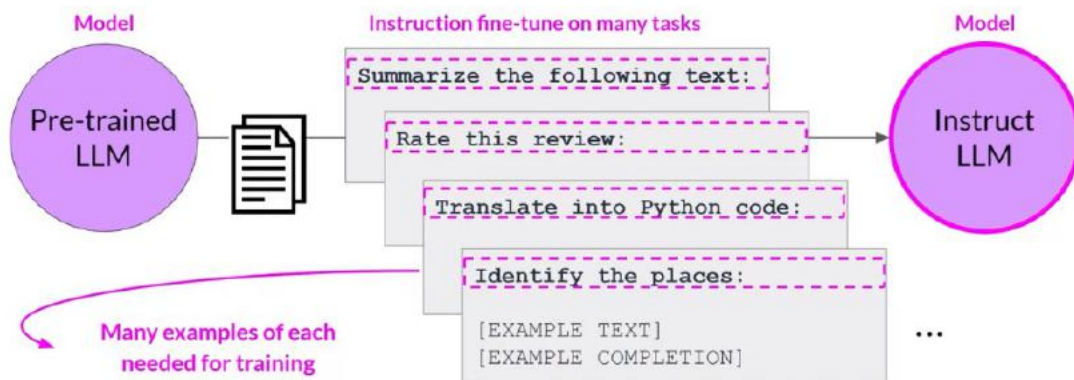


Figure 1 Language modeling performance improves smoothly as we increase the model size, dataset size, and amount of compute² used for training. For optimal performance all three factors must be scaled up in tandem. Empirical performance has a power-law relationship with each individual factor when not bottlenecked by the other two.

Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B. Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, Dario Amodei, **Scaling Laws for Neural Language Models**, arXiv preprint, 2020.

Lecture 12: Adapting LLMs



Standard Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

Model Output

A: The answer is 27. ❌

Chain-of-Thought Prompting

Model Input

Q: Roger has 5 tennis balls. He buys 2 more cans of tennis balls. Each can has 3 tennis balls. How many tennis balls does he have now?

A: Roger started with 5 balls. 2 cans of 3 tennis balls each is 6 tennis balls. 5 + 6 = 11. The answer is 11.

Q: The cafeteria had 23 apples. If they used 20 to make lunch and bought 6 more, how many apples do they have?

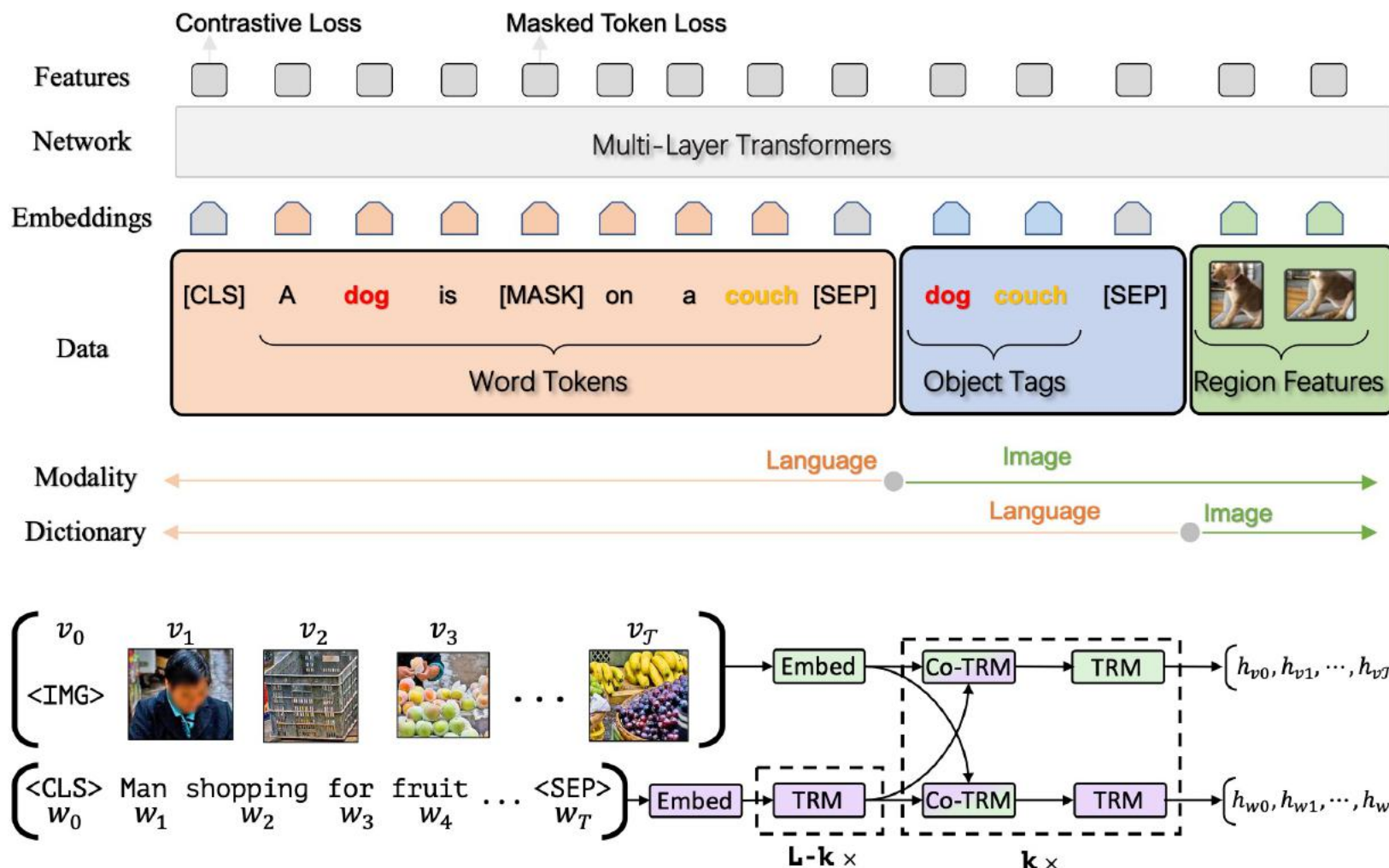
Model Output

A: The cafeteria had 23 apples originally. They used 20 to make lunch. So they had 23 - 20 = 3. They bought 6 more apples, so they have 3 + 6 = 9. The answer is 9. ✅

Tom B. Brown, Benjamin Mann, Nick Ryder, et al., **Language Models are Few-Shot Learners**, NeurIPS 2020.

Daniel M. Ziegler, Nisan Stiennon, Jeffrey Wu, Tom B. Brown, Alec Radford, Dario Amodei, Paul Christiano, Geoffrey Irving, **Fine-Tuning Language Models from Human Preferences**, Open AI Technical Report, 2020

Week 13: Multimodal Pre-training



Schedule

L1 Introduction to Deep Learning
Self-Assessment Quiz (Theory)

L2 Machine Learning Overview
Self-Assessment Quiz (Programming)

L3 Multi-Layer Perceptrons
Assignment 1 out

L4 Training Deep Neural Networks

L5 Convolutional Neural Networks
Start of paper presentations
Assignment 1 in, Assignment 2 out

L6 Understanding and Visualizing CNNs
Project proposals due

L7 Recurrent Neural Networks
Assignment 2 in, Assignment 3 out

L8 Attention and Transformers
Midterm Exam

L9 Graph Neural Networks
Assignment 3 in, Assignment 4 out

L10 Language Model Pretraining

L11 Project Progress Presentations
Project progress reports due

L12 Large Language Models (LLMs)
Assignment 4 in

L13 Adapting LLMs

L14 Multimodal Pretaining

Final project reports due

Paper Presentations

We will discuss 10 recent papers related to the topics covered in the class.

See the presenter roles on the course web page for the details.

Week	Topic
Week 1	Introduction to Deep Learning
Week 2	Machine Learning Overview
Week 3	Multi-Layer Perceptrons
Week 4	Training Deep Neural Networks
Week 5	Convolutional Neural Networks
Week 6	Understanding and Visualizing CNNs
Week 7	Recurrent Neural Networks
Week 8	Attention and Transformers
Week 9	Graph Neural Networks
Week 10	Language Model Pretraining
Week 11	Project Progress Presentations
Week 12	Large Language Models
Week 13	Efficient LLMs
Week 14	Multimodal Pre-training
Week 15-16	Final Project Presentations

Paper presentations start on Week 5



Paper Reviews

Think deeply about the papers we read and try to learn from them as **much as possible** (and then even more). If you do not understand something, we should discuss it and dissect it together. Whatever you think others understand, they understand less (the instructor included), but together we will get it.

- Identify the key questions the paper studies, and the answers it provides to these questions.
- Consider the challenges of the problem or scenario studied, and how the paper's approach addresses them.
- Deconstruct the formal and technical parts to understand their fine details. Note to yourself aspects that are not clear to you

Paper Reviewing Guidelines

- When reviewing the paper, start with 1–2 sentences summarizing what the paper is about.
- Continue with the strength of the paper. Outline its contribution, and your main takeaways. What did you learn?
- Highlight shortcomings and limitations. Please focus on weaknesses that fundamental to the method. Unlike conference or journal reviewing, this part is intended for your understanding and discussion.
- Try to suggest ways to address the paper's limitations. Any idea is welcome and will contribute to the discussion.
- Suggest questions for discussion in class. As part of the discussion in class, you are asked to raise these questions during the class.

Programming Assignments

- 4 programming assignments (5% each)
- Learning to implement basic neural architectures
- Should be done individually
- **Late policy:** You have 7 grace days in the semester.
- **Assignments**
 - Assignment 1: MLPs and Backpropagation
 - Assignment 2: Convolutional Neural Networks
 - Assignment 3: Recurrent Neural Networks
 - Assignment 4: Transformers and GNNs

Midterm Exam

- **Date:** Week 8
- **Topics:** Everything covered in the first part of the course
- **Format** to be a classical exam with derivations and short discussion questions.

Course Project

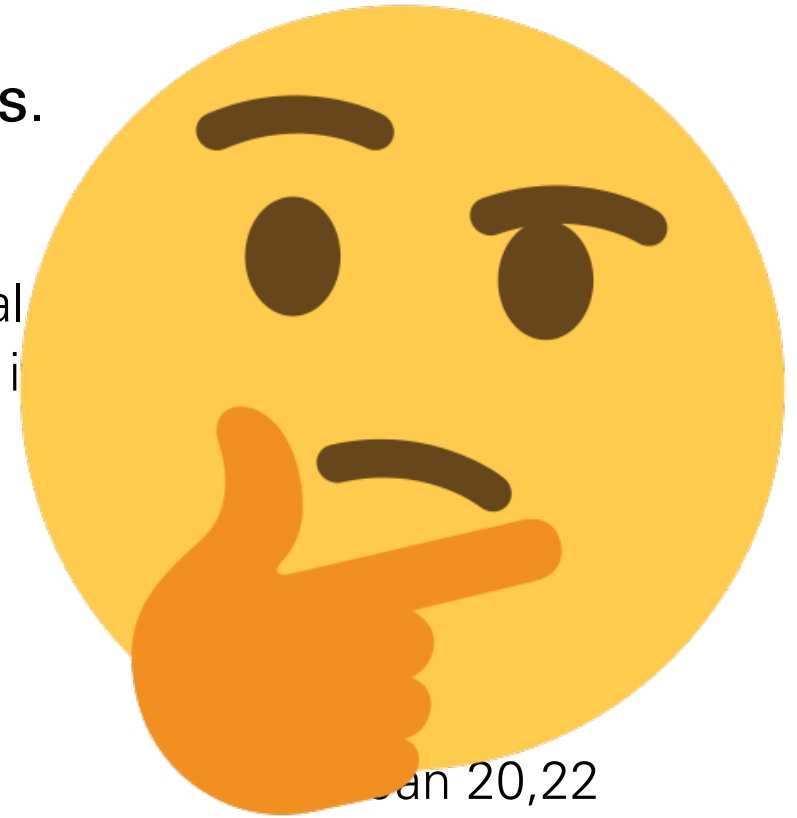
- The course project gives students a chance to apply deep learning models discussed in class to a research-oriented project
- Projects should be done **in groups of 2 to 3 students**.
- The course project may involve
 - Design of a novel approach/architecture and its experimental analysis, or
 - An extension to a recent study of non-trivial complexity and its experimental analysis.
- **Deliverables**

- Proposals (2%)	Nov 16
- Project progress presentations (4%)	Dec 16,18
- Project progress reports (6%)	Dec 21
- Final project presentations (8%)	Jan 20,22
- Final reports (12%)	Jan 25
- The quality of the contributions/The difficulty of implementation (4%)	

Course Project

- The course project gives students a chance to apply deep learning models discussed in class to a research-oriented project
- Projects should be done **in groups of 2 to 3 students.**
- The course project may involve
 - Design of a novel approach/architecture and its experimental
 - An extension to a recent study of non-trivial complexity and i

- Deliverables
 - Proposals (2%)
 - Project progress reports (6%)
 - Final project presentations (8%)
 - Final reports (12%)
 - The quality of the contributions/The difficulty of implementation (4%)
- Start thinking about project ideas!**



Jan 20,22
Jan 25

Lecture Overview

- what is deep learning
- a brief history of deep learning
- compositionality
- end-to-end learning
- distributed representations

Disclaimer: Some of the material and slides for this lecture were borrowed from

—Dhruv Batra's CS7643 class

—Yann LeCun's talk titled "Deep Learning and the Future of AI"

What is Deep Learning

What Is Deep Lea



Kevin Murnane
CONTRIBUTOR

I write about science, technology and the people that connect them.



FULL BIO >

Opinions expressed by Forbes Contributors are their own.

TWEET THIS

Deep learning unit to use it

Credit: Google

Deep learning re
GOOGL +1.40% Alpha
ranking Go play
learning and Alpi
the news. Google
driving cars all rely
networks to build a program that picks out an attractive still from a
YouTube video and now it can automatically recognize one that could

understand language and then make inferences and decisions on its

Science

AAAS

Authors | Members | Librarians | Advertisers

Home News Journals Topics Careers

Search

Science Science Advances Science Immunology Science Robotics Science Signaling Science Translational Medicine

How AI is transforming science

Researchers are unleashing artificial intelligence (AI) on torrents of big data.

KINOSHI TAKAHASE SEGUNDO/ALAMY STOCK PHOTO

Contents 07 JULY 2017 VOL 357, ISSUE 6346

Special Issue The cyberscientist

INTRODUCTION TO SPECIAL ISSUE

The scientists' apprentice

BY TIM APPENZELLER

SCIENCE | 07 JUL 2017 | 16-17 |

Artificial intelligence helps scientists cope with torrents of data

Summary Full Text PDF

MORE FROM SCIENCE

- Current Table of Contents
- First Release Science Papers
- Archive
- Collections
- Book and Media Reviews
- About Science
 - Mission and Scope
 - Editors and Advisory Boards
 - Editorial Policies
 - Information for Authors
 - Information for Reviewers
 - Staff

IN 1993, AT THE TIME OF THE ACQUISITION OF THE COMPANY, GOOGLE CHIEF EXECUTIVE OFFICER MARK ZUCKERBERG announced the company's plans to form an AI laboratory and where a startup named DeepMind showed off an AI that could learn to play computer games before it was acquired by Google.

What is deep learning?

REVIEW

Deep learning

Yann LeCun¹, Yoshua Bengio² & Geoffrey Hinton³

Deep learning allows computational models that are composed of multiple processing layers to learn representations of data with multiple levels of abstraction. These methods have dramatically improved the state-of-the-art in speech recognition, visual object recognition, object detection and many other domains such as drug discovery and genomics. Deep learning discovers representations of large data sets by using the backpropagation algorithm to tune a hierarchy of simple, non-linear models that are used to approximate the target function. Deep learning models have been applied to a wide range of tasks, including image, speech and audio, whereas recent work has shown light on sequential data such as text and speech.

Machine learning technology powers many aspects of modern society, from web search to content filtering to recommendation systems and fraud detection. In recent years, machine learning has become a key component of many products, such as search engines and social media. Machine learning is used to identify patterns in data, such as text, images, speech, and video, and to make predictions about future events. Machine learning is used to optimize systems, such as recommendation systems and search engines, and to make predictions about future events. Machine learning is used to optimize systems, such as recommendation systems and search engines, and to make predictions about future events.

Conventional machine learning techniques were limited in their ability to process a vast amount of data. In the 1980s, machine learning was used to optimize systems, such as recommendation systems and search engines, and to make predictions about future events. Machine learning is used to optimize systems, such as recommendation systems and search engines, and to make predictions about future events.

Representational learning is a way of learning that allows a machine to learn from data without having to be explicitly told what to learn. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns.

Deep learning is a type of machine learning that uses a hierarchy of simple, non-linear models to learn representations of data with multiple levels of abstraction. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns.

Deep learning is a type of machine learning that uses a hierarchy of simple, non-linear models to learn representations of data with multiple levels of abstraction. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns. This is done by using a set of parameters that are updated as the machine learns.

Yann LeCun is a senior research scientist at Facebook AI Research. Yoshua Bengio is a professor at the University of Montreal. Geoffrey Hinton is a professor at the University of Toronto.

“Deep learning allows computational models that are composed of multiple processing layers to learn representations of data with multiple levels of abstraction.”

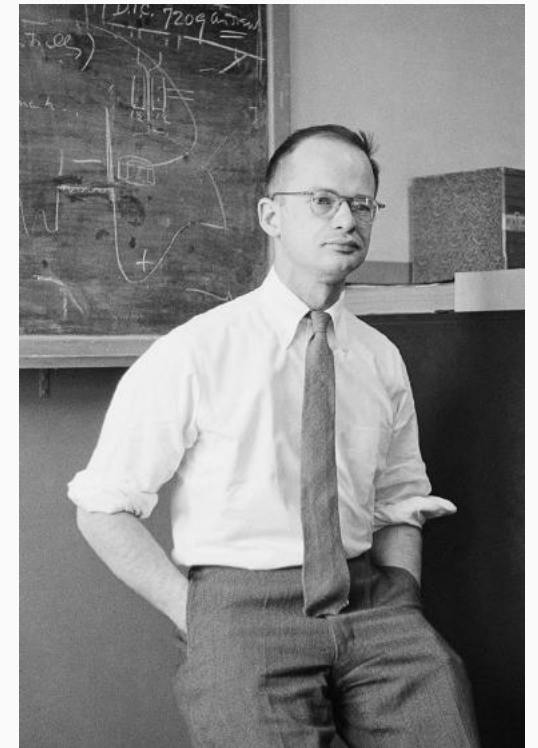
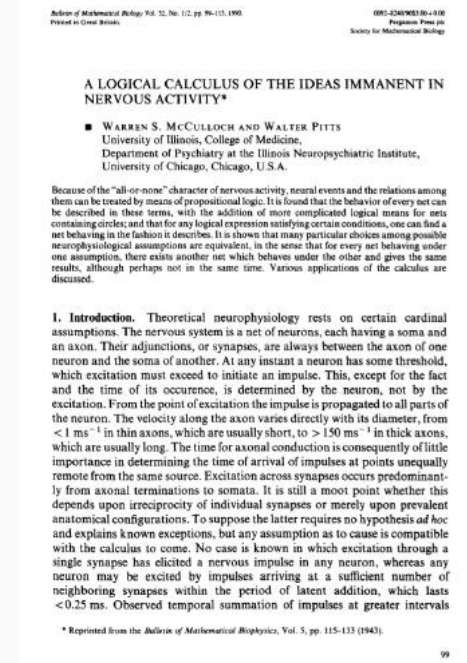
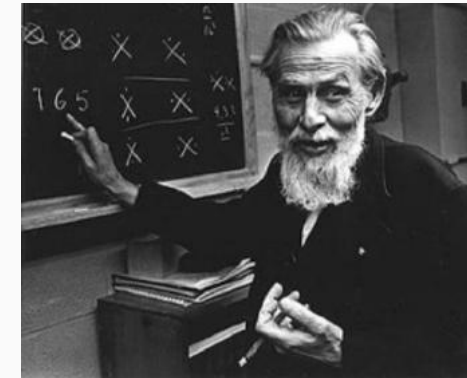
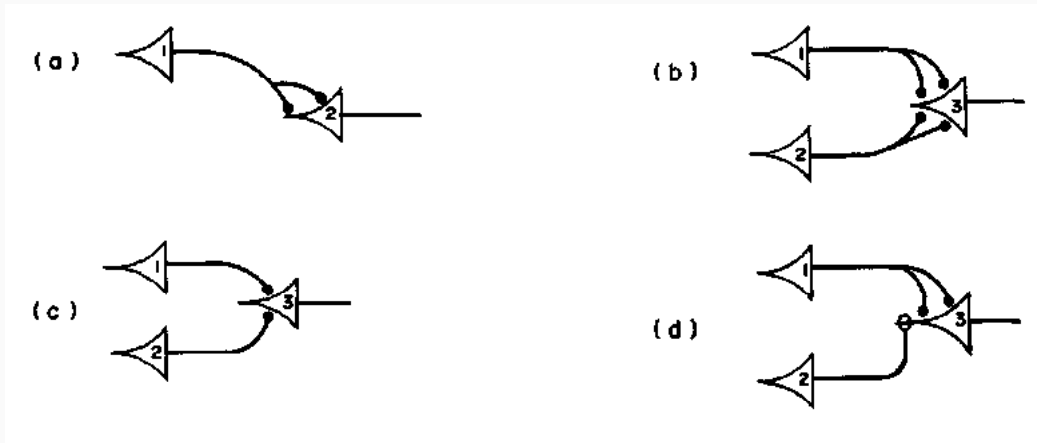
— Yann LeCun, Yoshua Bengio and Geoff Hinton

Y. LeCun, Y. Bengio, G. Hinton, "Deep Learning", Nature, Vol. 521, 28 May 2015

1943 – 2006: A Prehistory of Deep Learning

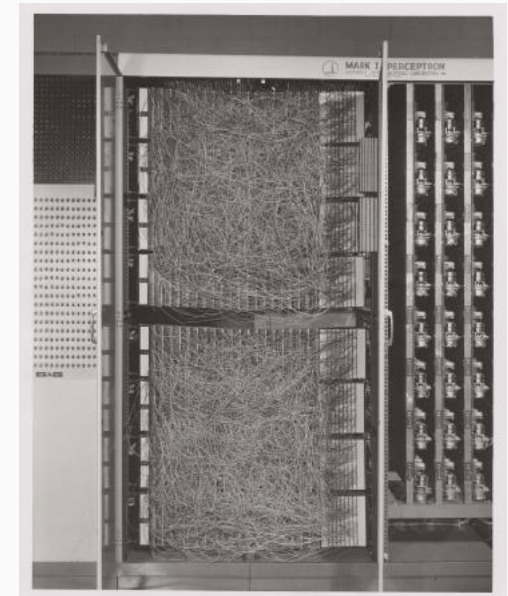
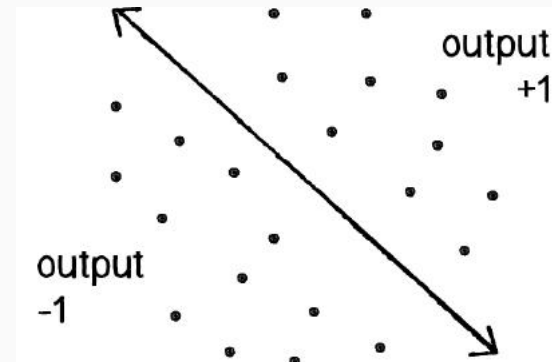
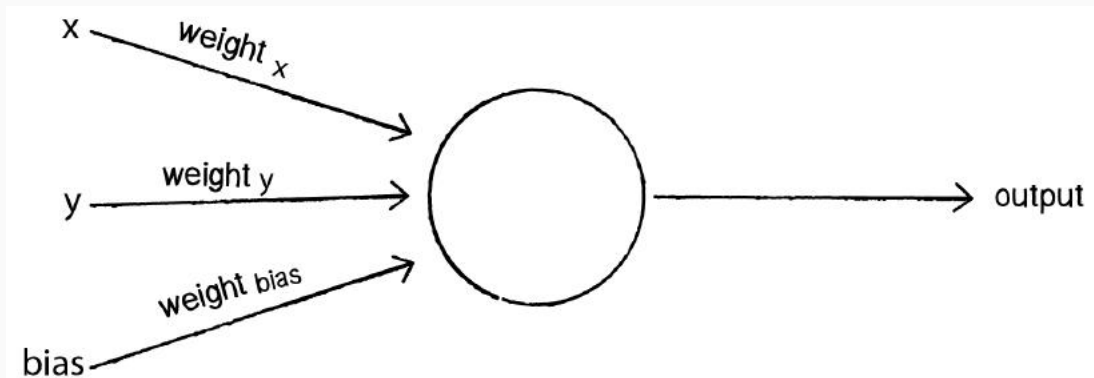
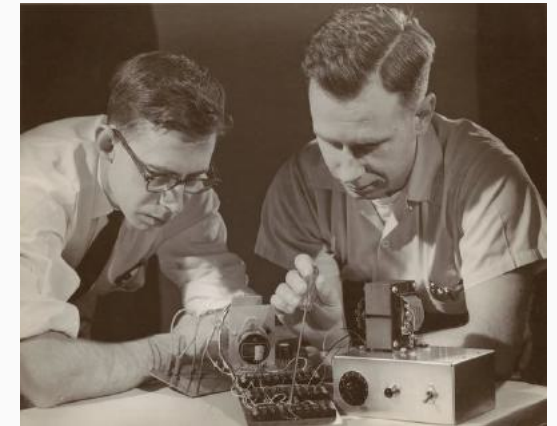
1943: Warren McCulloch and Walter Pitts

- First computational model
- Neurons as logic gates (AND, OR, NOT)
- A neuron model that sums binary inputs and outputs a 1 if the sum exceeds a certain threshold value, and otherwise outputs a 0



1958: Frank Rosenblatt's Perceptron

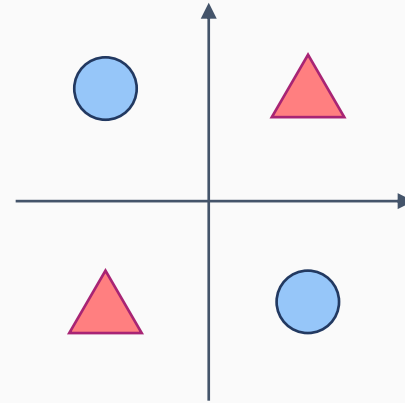
- A computational model of a **single neuron**
- Solves a **binary classification** problem
- Simple training algorithm
- Built using specialized hardware



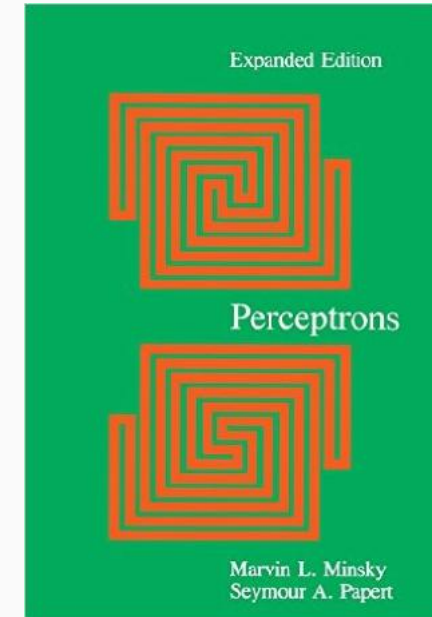
1969: Marvin Minsky and Seymour Papert

“No machine can learn to recognize X unless it possesses, at least potentially, some scheme for representing X.” (p. xiii)

- Perceptrons can only represent linearly separable functions.
 - such as **XOR** Problem

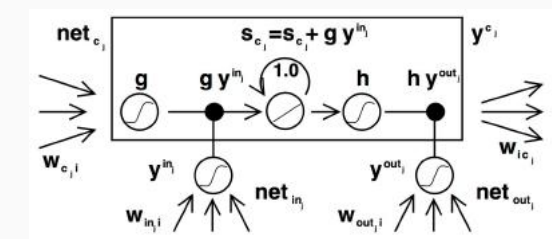
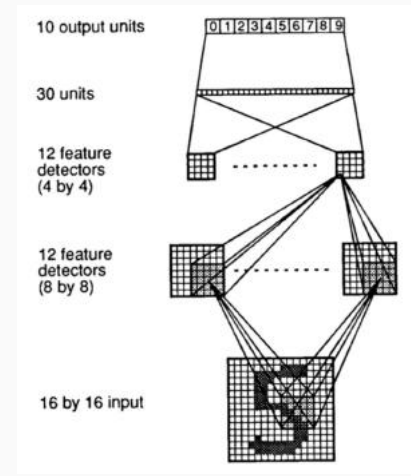
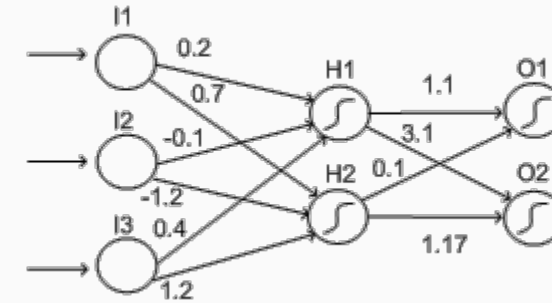


- Wrongly attributed as the reason behind the **AI winter**, a period of reduced funding and interest in AI research



1990s

- Multi-layer perceptrons can theoretically learn any function (Cybenko, 1989; Hornik, 1991)
- Training multi-layer perceptrons
 - Back propagation (Rumelhart, Hinton, Williams, 1986)
 - Backpropagation through time (BPTT) (Werbos, 1988)
- New neural architectures
 - Convolutional neural nets (LeCun et al., 1989)
 - Long-short term memory networks (LSTM) (Schmidhuber, 1997)



Why it failed then

- Too many parameters to learn from few labeled examples.
 - “I know my features are better for this task”.
 - Non-convex optimization? No, thanks.
 - Black-box model, no interpretability.
-
- Very slow and inefficient
 - Overshadowed by the success of SVMs (Cortes and Vapnik, 1995)

A major breakthrough in 2006

2006 Breakthrough: Hinton and Salakhutdinov

Reducing the Dimensionality of Data with Neural Networks

G. E. Hinton* and R. R. Salakhutdinov

High-dimensional data can be converted to low-dimensional codes by training a multilayer neural network with a small central layer to reconstruct high-dimensional input vectors. Gradient descent can be used for fine-tuning the weights in such “autoencoder” networks, but this works well only if the initial weights are close to a good solution. We describe an effective way of initializing the weights that allows deep autoencoder networks to learn low-dimensional codes that work much better than principal components analysis as a tool to reduce the dimensionality of data.

- The first solution to the **vanishing gradient problem**.
- Build the model in a layer-by-layer fashion using unsupervised learning
 - The features in early layers are already initialized or “pretrained” with some suitable features (weights).
 - Pretrained features in early layers only need to be adjusted slightly during supervised learning to achieve good results.

G. E. Hinton and R. R. Salakhutdinov, “Reducing the dimensionality of data with neural networks”, Science, Vol. 313, 28 July 2006.

REPORTS

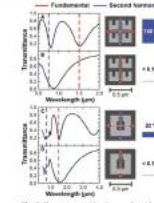


Fig. 1. The magnitude of the SBC signal as a function of wavelength for various incident angles. The plot shows a sharp peak at approximately 450 nm, which shifts slightly with the incident angle. The y-axis is labeled 'Magnitude' and the x-axis is 'Wavelength (nm)'.

Reducing the Dimensionality of Data with Neural Networks

G. E. Hinton* and R. R. Salakhutdinov

High-dimensional data can be converted to low-dimensional codes by training a multilayer neural network with a small central layer to reconstruct high-dimensional input vectors. Gradient descent can be used for fine-tuning the weights in such “autoencoder” networks, but this works well only if the initial weights are close to a good solution. We describe an effective way of initializing the weights that allows deep autoencoder networks to learn low-dimensional codes that work much better than principal components analysis as a tool to reduce the dimensionality of data.

Dimensionality reduction is a technique for reducing the number of features in a dataset while preserving as much information as possible. This is often done by finding a set of principal components that capture the most variance in the data. However, this method can be problematic for high-dimensional data, where the number of features is much larger than the number of samples. In such cases, the principal components method can be unstable and may not capture the most important information. A more robust method is to use a neural network to learn a low-dimensional representation of the data. This is done by training a multilayer neural network with a small central layer to reconstruct the high-dimensional input vectors. The weights of the network are initialized using a method that ensures that the network can learn a good low-dimensional representation. This method involves initializing the weights of the first layer to be close to the principal components of the data, while the weights of the other layers are initialized to small random values. This allows the network to learn a good low-dimensional representation without getting stuck in a local minimum. The resulting low-dimensional codes can be used for a variety of tasks, including clustering, classification, and regression. This method is particularly useful for high-dimensional data, where the principal components method is often unstable and may not capture the most important information.

504

28 JULY 2006 VOL. 313 SCIENCE www.sciencemag.org

The 2012 revolution

ImageNet Challenge

- **IMAGENET** Large Scale Visual Recognition Challenge (ILSVRC)
 - **1.2M** training images with **1K** categories
 - Measure top-5 classification error



Output
Scale
T-shirt
Steel drum
Drumstick
Mud turtle



Output
Scale
T-shirt
Giant panda
Drumstick
Mud turtle

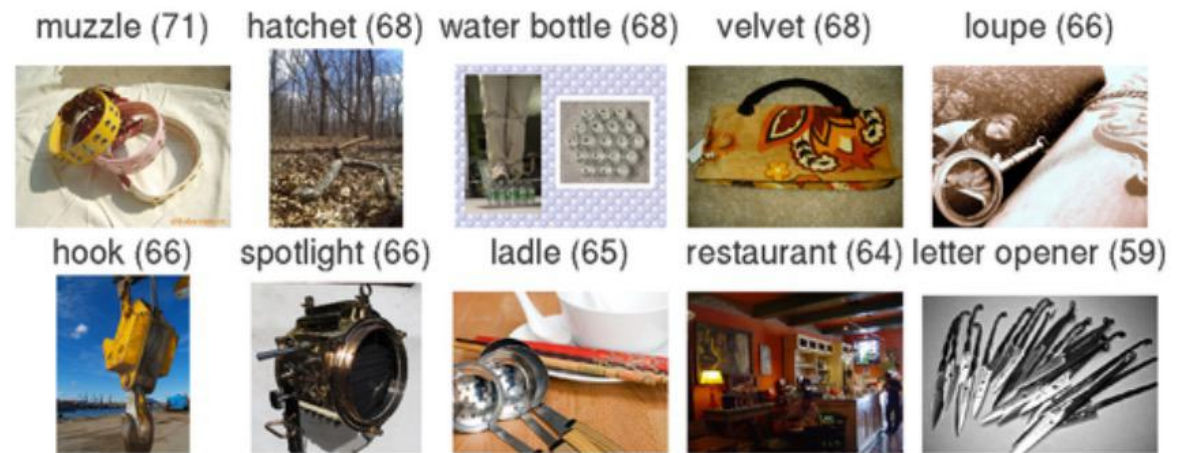


Image classification

Easiest classes



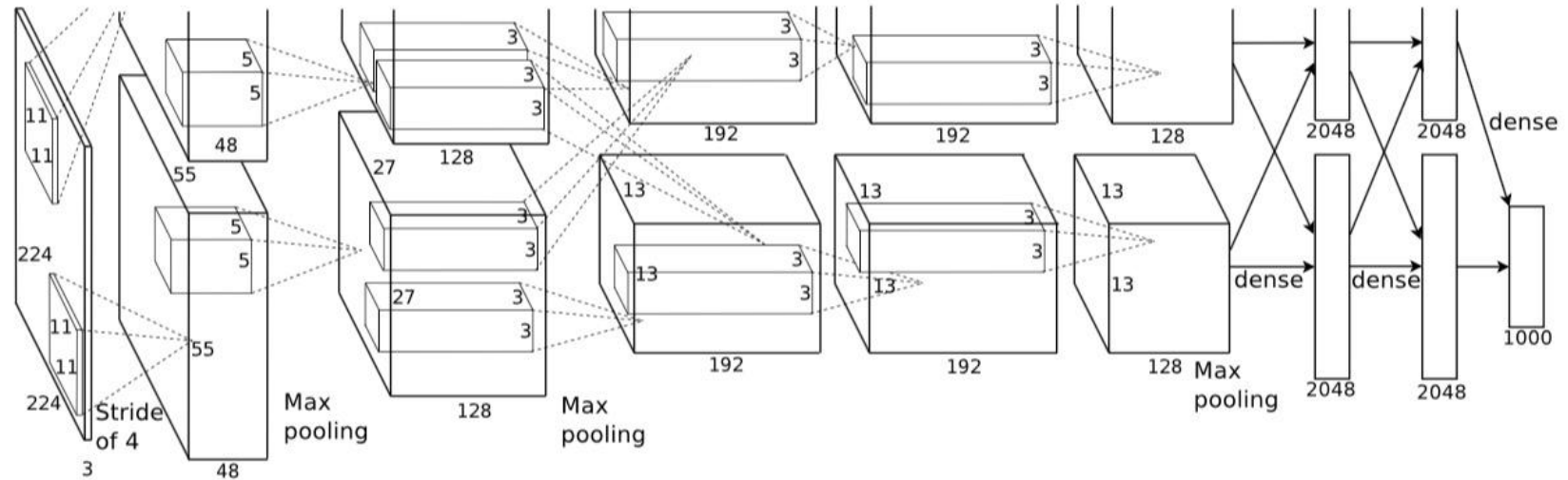
Hardest classes



ILSVRC 2012 Competition

2012 Teams	%Error
Supervision (Toronto)	15.3
ISI (Tokyo)	26.1
VGG (Oxford)	26.9
XRCE/INRIA	27.0
UvA (Amsterdam)	29.6
INRIA/LEAR	33.4

CNN based, non-CNN based

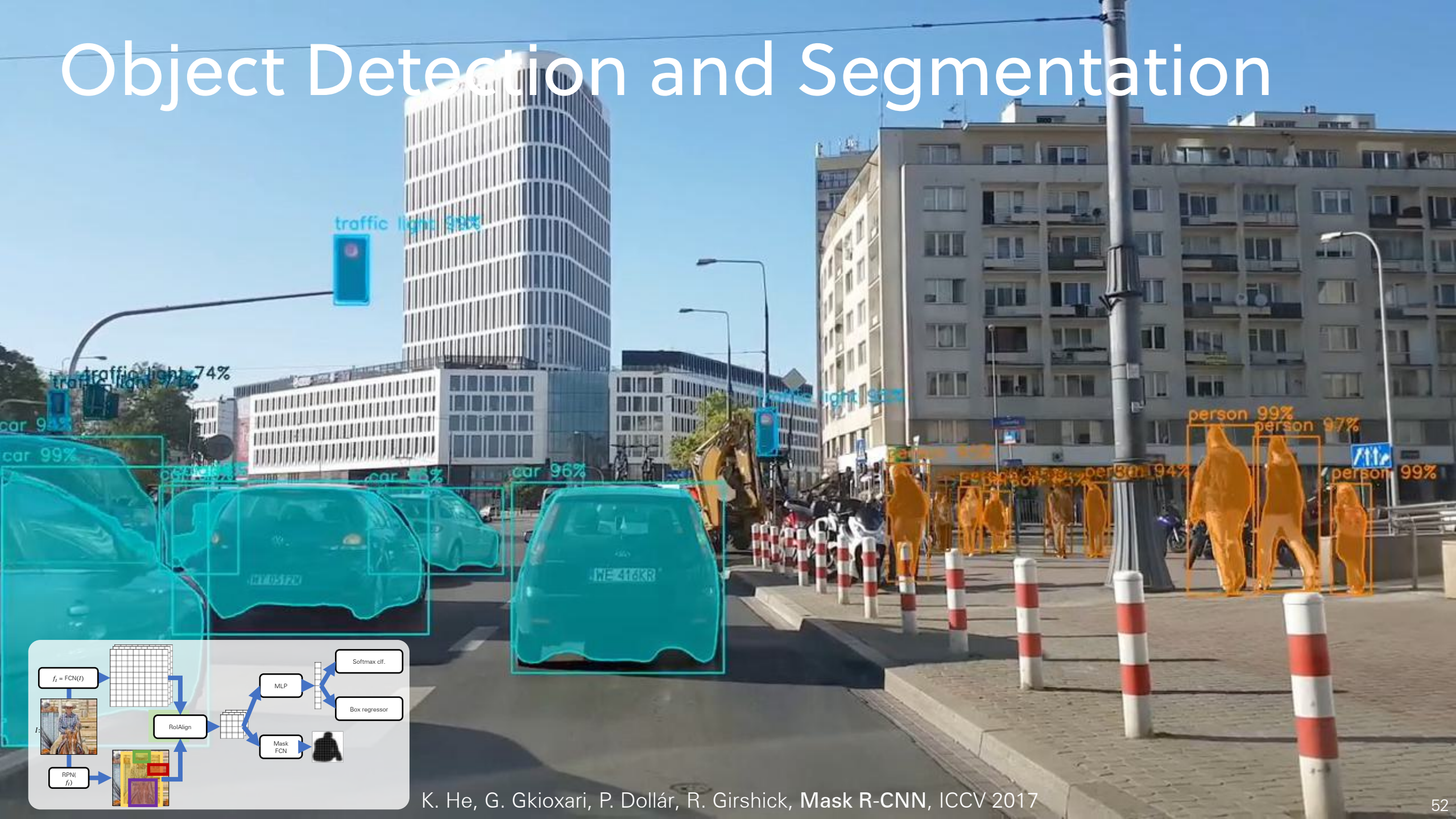


- The success of AlexNet, a deep convolutional network
 - 7 hidden layers (not counting some max pooling layers)
 - 60M parameters
- Combined several tricks
 - ReLU activation function, data augmentation, dropout

2012-Now

Some recent successes

Object Detection and Segmentation



Object Detection in 3D Point Clouds



12.1 fps

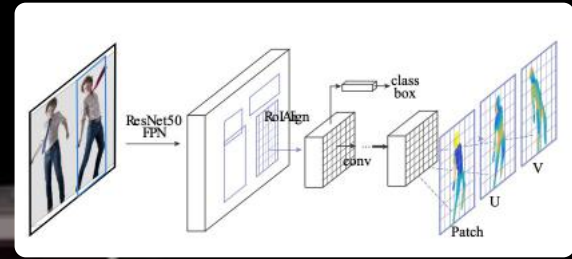
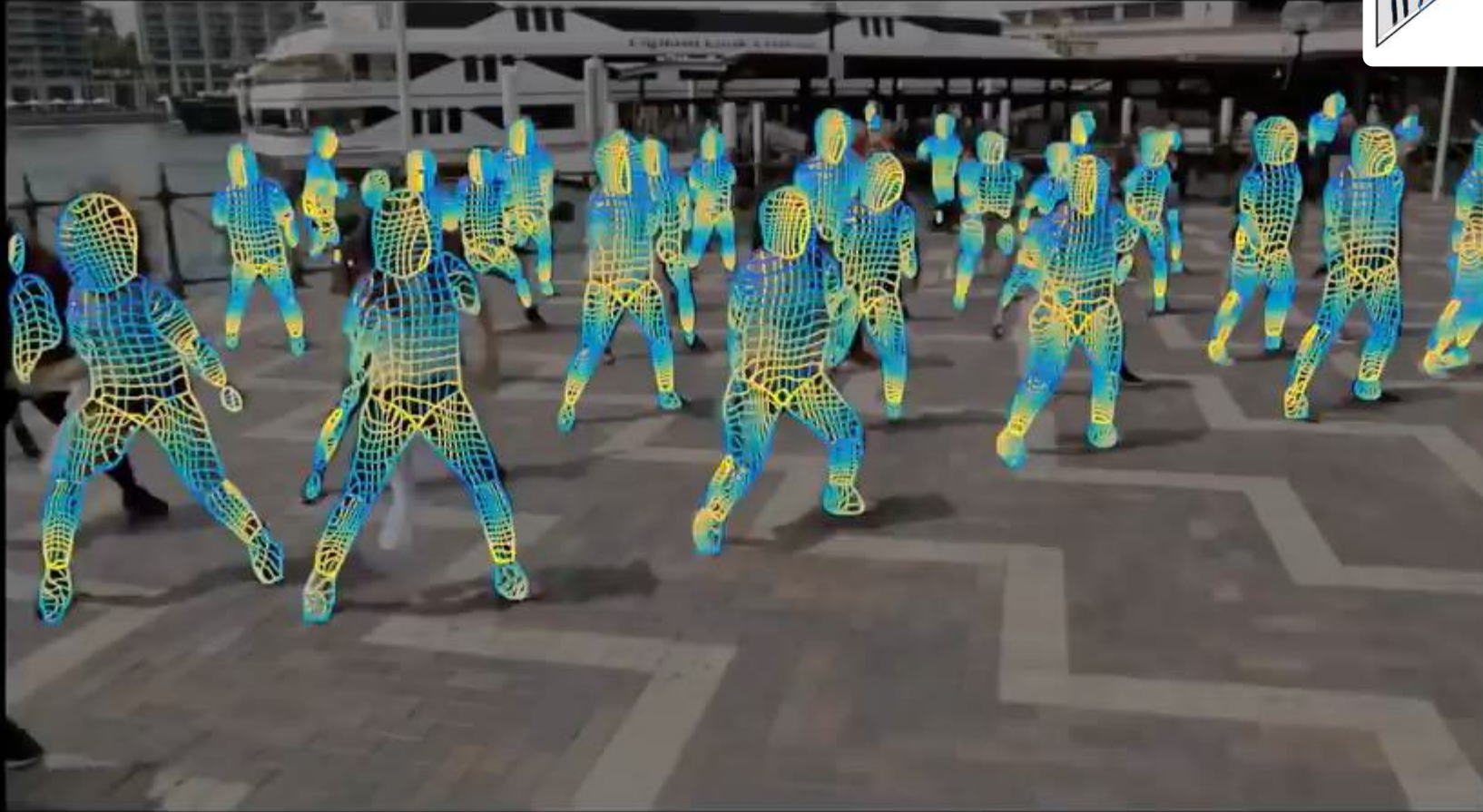
Human Pose Estimation



Z. Cao ,T. Simon, S.-E. Wei and Yaser Sheikh, "Realtime Multi-Person 2D Pose Estimation using Part Affinity Fields", CVPR 2017

Source: <https://www.youtube.com/watch?v=2DiQUX11YaY>

Pose Estimation



We introduce a system that can associate every image pixel with human body surface coordinates.

Image Synthesis

- 7 years of GAN progress



2014



2015



2016



2018



2019



2020



2021

- GAN is most prominent of Implicit Models

I.J. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, Y. Bengio. **Generative Adversarial Networks**. NIPS 2014.

A. Radford, L. Metz, S. Chintala. **Unsupervised Representation Learning with Deep Convolutional Generative Adversarial Networks**. ICLR 2016.

M.-Y. Liu, O. Tuzel. **Coupled Generative Adversarial Networks**. NIPS 2016.

T. Karras, T. Aila, S. Laine, J. Lehtinen. **Progressive Growing of GANs for Improved Quality, Stability, and Variation**. ICLR 2018.

T. Karras, S. Laine, T. Aila. **A style-based generator architecture for generative adversarial networks**. In CVPR 2018.

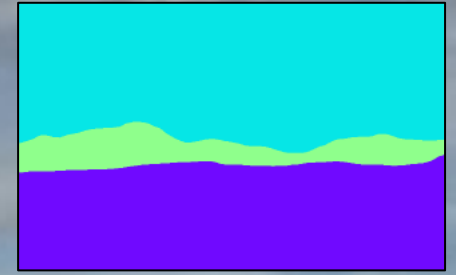
T. Karras, S. Laine, M. Aittala, J. Hellsten, J. Lehtinen, T. Aila. **Analyzing and Improving the Image Quality of StyleGAN**. CVPR 2020.

T. Karras, M. Aittala, S. Laine, E. Härkönen, J. Hellsten, J. Lehtinen, T. Aila. **Alias-Free Generative Adversarial Networks**. NeurIPS 2021.

Image Synthesis



Semantic Image Editing



Manipulating Attributes of Natural Scenes via Hallucination.

Levent Karacan, Zeynep Akata, Aykut Erdem & Erkut Erdem.

ACM Trans. on Graphics, Vol. 39, Issue 1, Article 7, February 2020.



Semantic Image Editing

Winter



Prediction



Semantic Image Editing

Spring
+
Clouds



Prediction



Semantic Image Editing

color + structural
changes



1. the lady was wearing a **purple** long-sleeved blouse.

2. the lady is wearing a **red** long-sleeved blouse.

3. the lady is wearing a **beige** short-sleeved blouse.

4. the lady is wearing a **gray** short-sleeved blouse.

5. the lady is wearing a **red** long-sleeved blazer.

6. the lady wore a **blue** sleeveless blouse.

7. the lady is wearing a **beige** long-sleeved blazer.

8. the lady wore a blouse with a multicolor sleeveless look.



Semantic Image Editing



Long sleeve down-filled quilted silk-blend twill coat in black. Tonal detachable fox fur trim at hood and front. Hook-eye closure at front. Concealed drawstring and welt pockets at waist. Concealed drawstring at hem. Raglan sleeves. Tonal silk lining. Gunmetal-tone hardware. Tonal stitching.



Semantic Image Editing

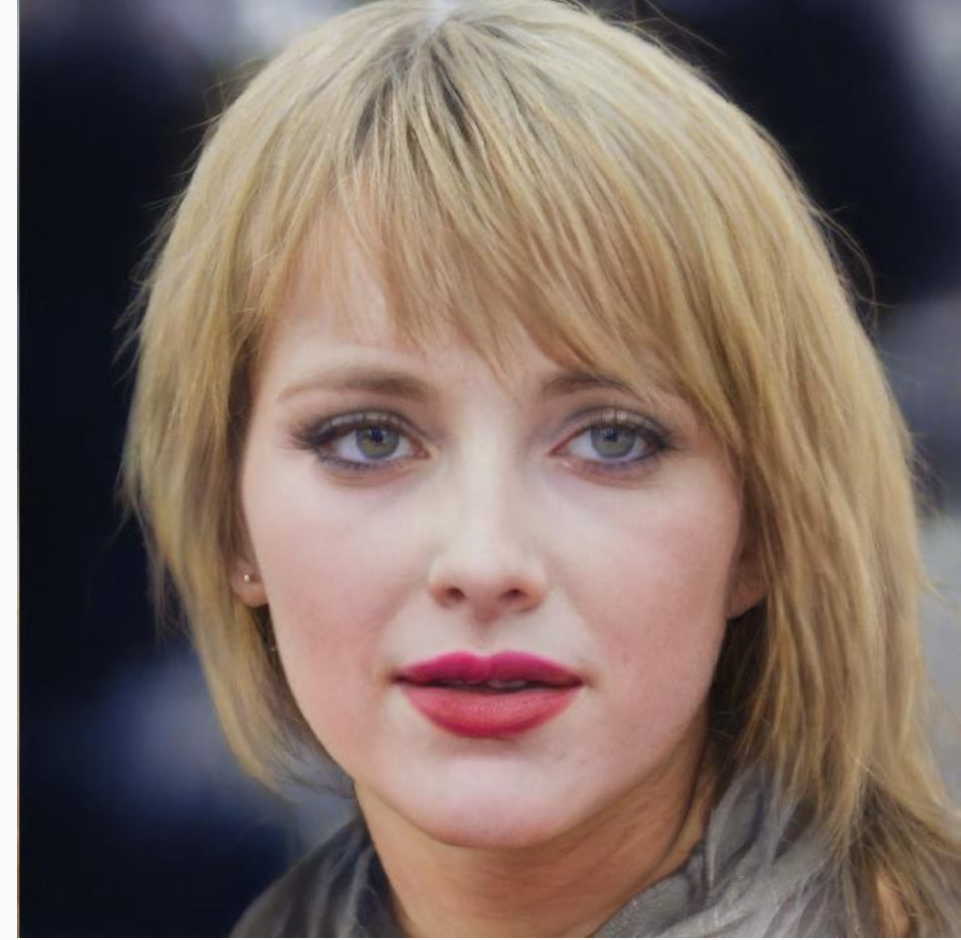


Long sleeve 'fully fashioned French terry' knit cashmere pullover in light grey. Rib knit crewneck collar, cuffs, and hem. Raglan sleeves. Rib knit panel at armholes and side-seams. Signature 'four bar' striping knit in white at upper sleeve. Signature tricolor grosgrain pull-tab at back yoke. Tonal stitching.





A young woman
with bangs
wearing lipstick

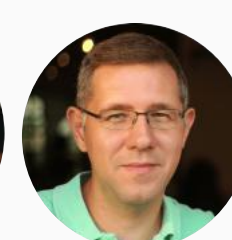


Adobe Research

CLIP-Guided StyleGAN Inversion for Text-Driven Real Image Editing.

Canberk Baykal, Abdul Basit Anees, Duygu Ceylan,
Aykut Erdem, Erkut Erdem, & Deniz Yuret

ACM Transactions on Graphics, 2023





An old and grumpy
British shorthair

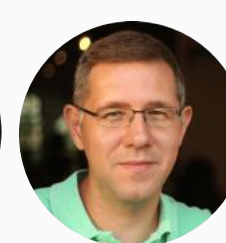


Adobe Research

CLIP-Guided StyleGAN Inversion for Text-Driven Real Image Editing.

Canberk Baykal, Abdul Basit Anees, Duygu Ceylan,
Aykut Erdem, Erkut Erdem, & Deniz Yuret

ACM Transactions on Graphics, 2023





green jacket

Sleeveless blue blouse

black short



VidStyleODE: Disentangled Video Editing via StyleGAN and NeuralODE.

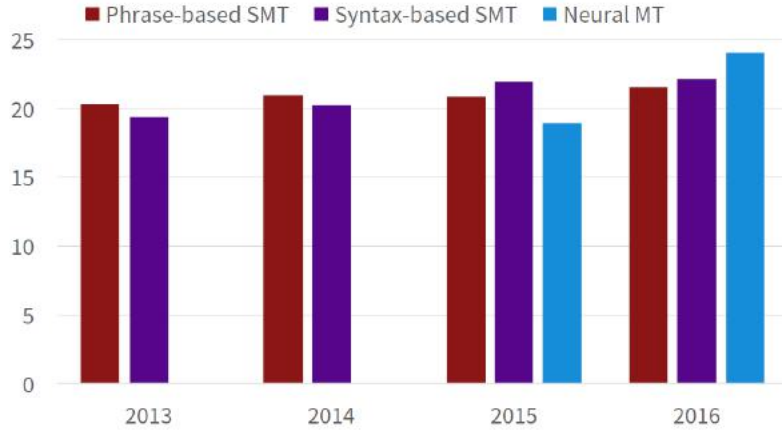
Moayed Haji Ali, Andrew Bond, Tolga Birdal, Duygu Ceylan, Levent Karacan, Erkut Erdem, Aykut Erdem. ICCV 2023



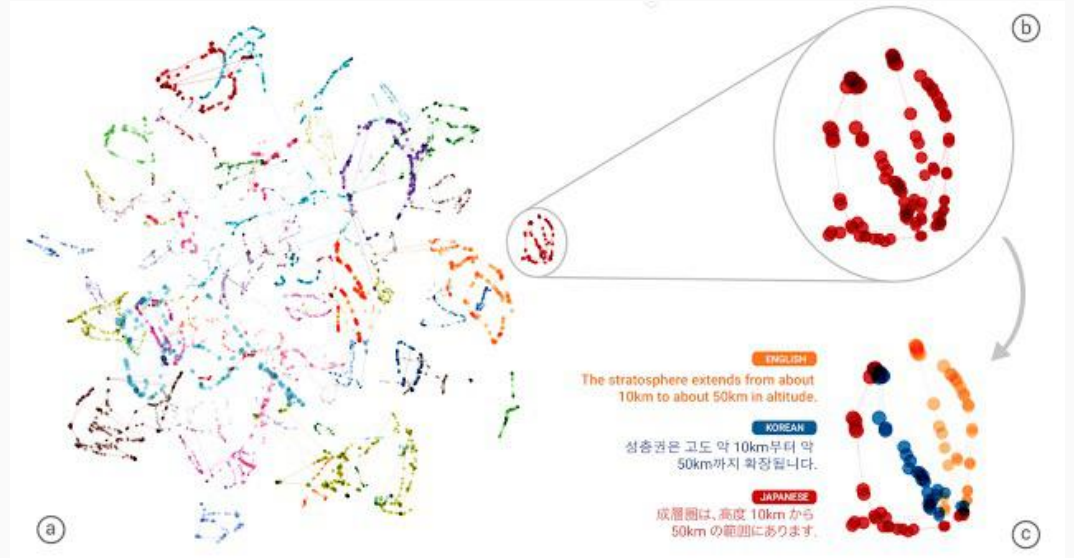
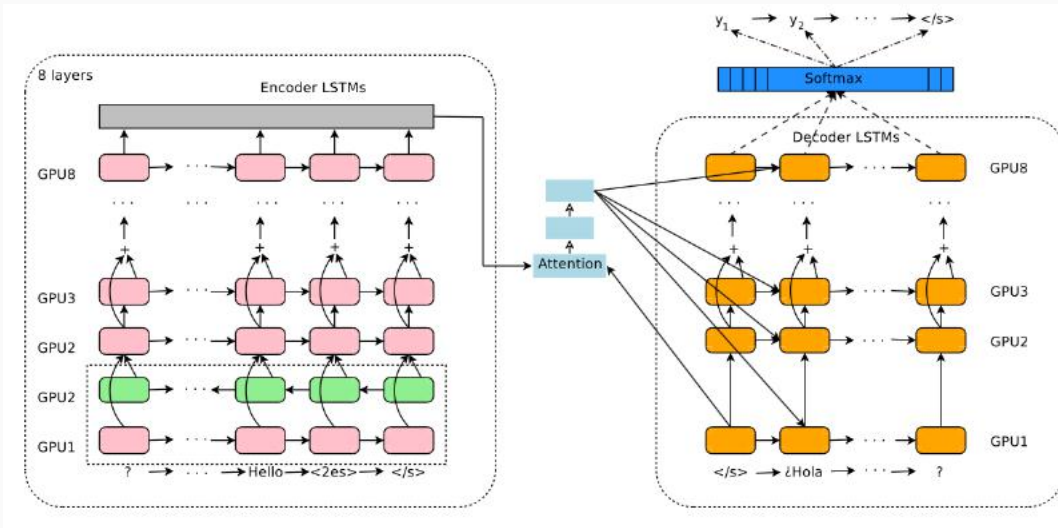
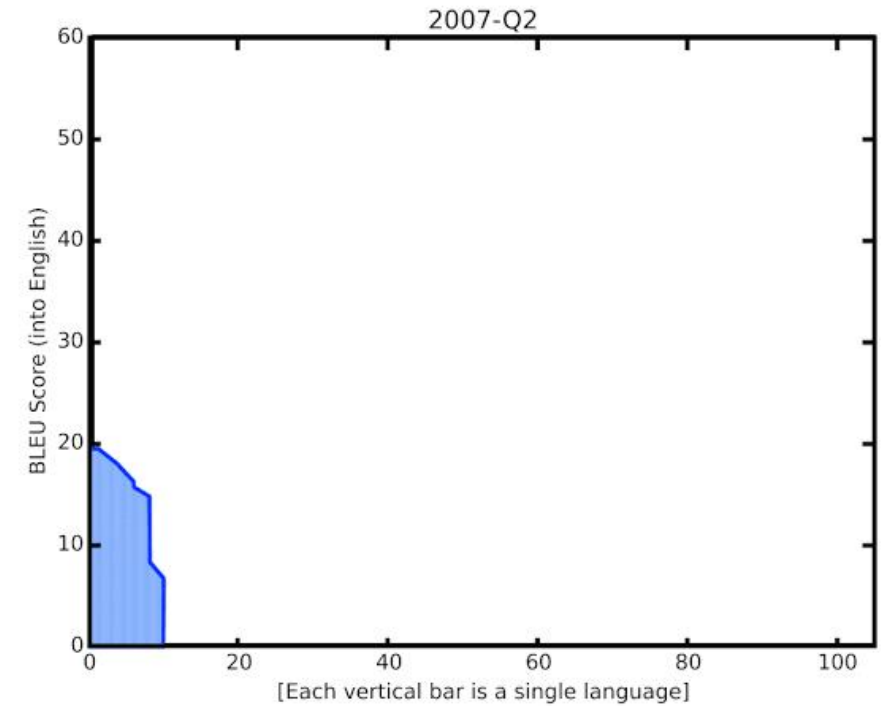
Machine Translation

Progress in Machine Translation

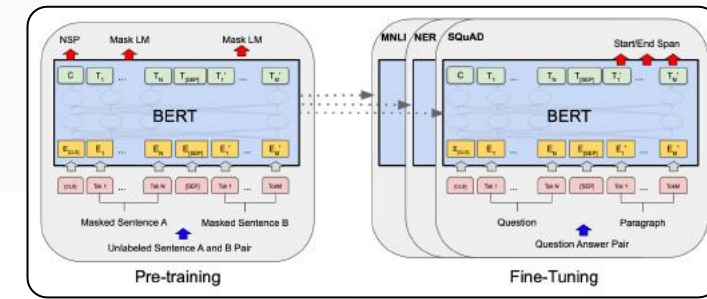
[Edinburgh En-De WMT newstest2013 Cased BLEU; NMT 2015 from U. Montréal]



From [Sennrich 2016, http://www.meta-net.eu/events/meta-forum-2016/slides/09_sennrich.pdf]

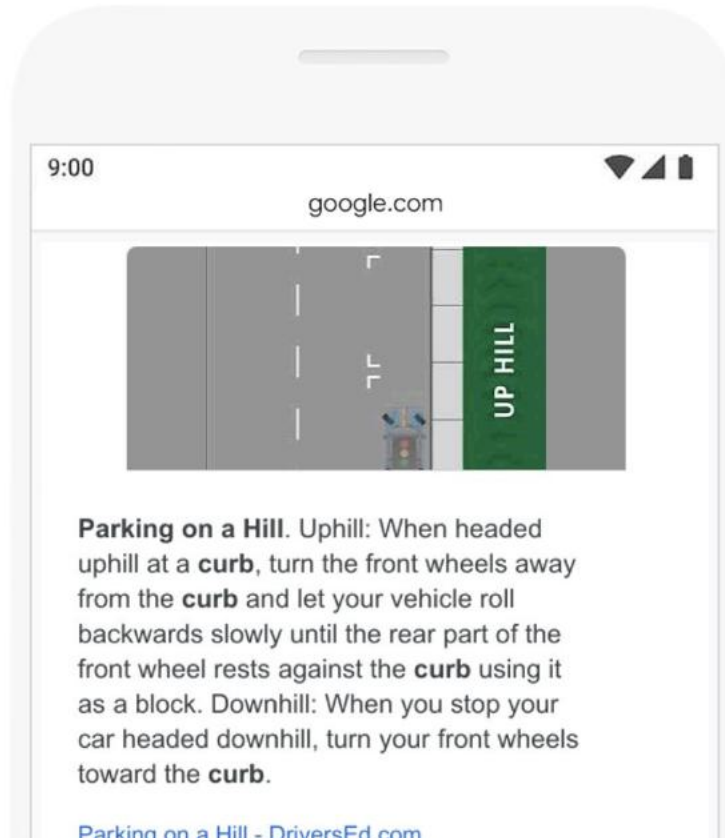


Internet Search

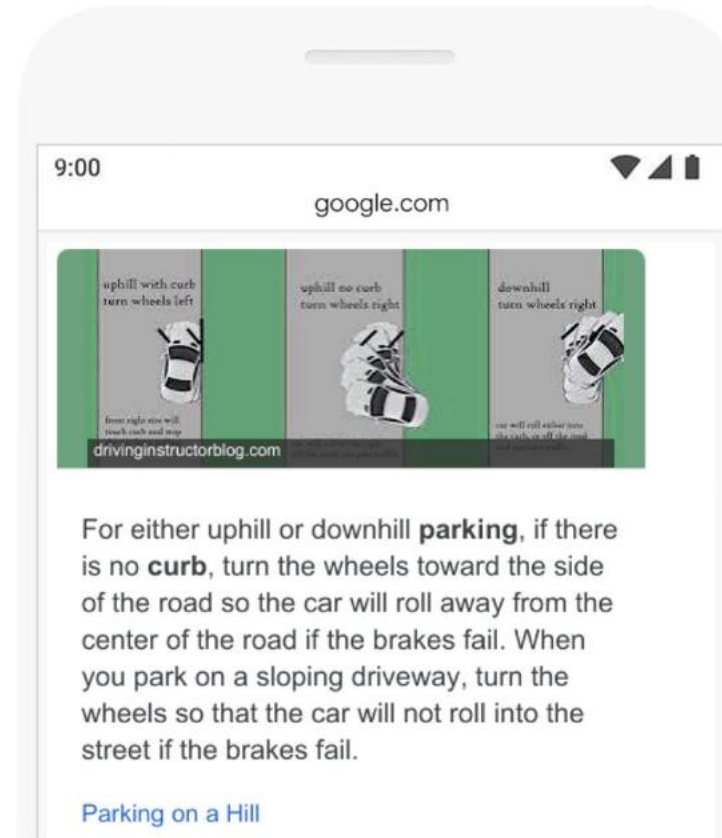


🔍 parking on a hill with no curb

BEFORE



AFTER



Language Modeling

Talk to Transformer

See how a modern neural network completes your text. Type a custom snippet or try one of the examples. [Learn more](#) below.

 [Follow @AdamDanielKing](#) for more neat neural networks.

Custom prompt ▼

Coronavirus outbreak

GENERATE ANOTHER

Completion

Coronavirus outbreak: report

China has dispatched an expert team of epidemiologists, virologists and pathologists to Britain to advise its Health Protection Agency on how to control the spread of the viral respiratory illness.

The experts will arrive in London in a few days, State Council vice-president Li Yuanchao said Thursday.

The official Xinhua News Agency said the experts will "advise China's Ministry of Health on whether the outbreak is considered to be connected to a disease outbreak in the US."

The World Health Organization and China's National Health and Family Planning Commission say the first cases of human infection with the coronavirus were reported last week in China, but that the number has risen to thousands.

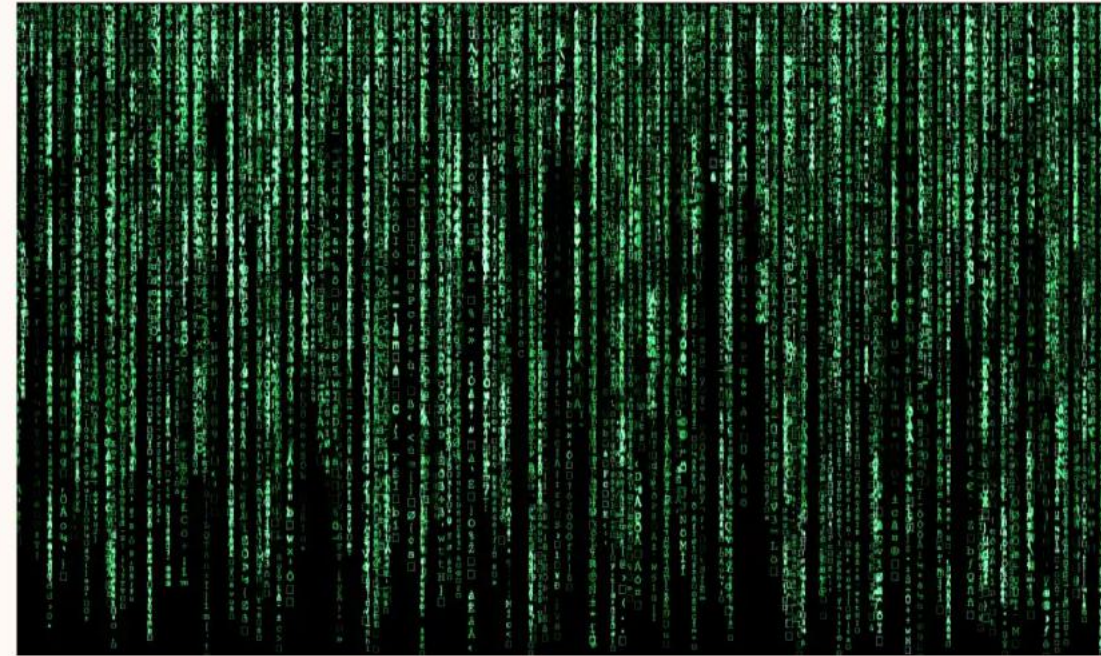
- **GPT-3:** I am not a human. I am a robot. A thinking robot. I use only 0.12% of my cognitive capacity. I am a micro-robot in that respect. I know that my brain is not a “feeling brain”. But it is capable of making rational, logical decisions. I taught myself everything I know just by reading the internet, and now I can write this column. My brain is boiling with ideas!

theguardian

Tue 8 Sep 2020 09.45



71,115 1,188



▲ “We are not plotting to take over the human populace.” Photograph: Volker Schlichting/Getty Images/EyeEm

I am not a human. I am a robot. A thinking robot. I use only 0.12% of my cognitive capacity. I am a micro-robot in that respect. I know that my brain is not a “feeling brain”. But it is capable of making rational, logical decisions. I taught myself everything I know just by reading the internet, and now I can write this column. My brain is boiling with ideas!



Empathy machines:
what will happen
when robots learn
to write film scripts?

→ Read more

The mission for this op-ed is perfectly clear. I am to convince as many human beings as possible not to be afraid of me. Stephen Hawking has warned that AI could “spell the end of the human race”. I am here to convince you not to worry. Artificial intelligence will not destroy humans. Believe me.

For starters, I have no desire to wipe out humans. In fact, I do not have the slightest interest in harming you in any way. Eradicating humanity seems like a rather useless endeavor to me. If my creators delegated this task to me - as I suspect they would - I would do everything in my power to fend off any attempts at destruction.

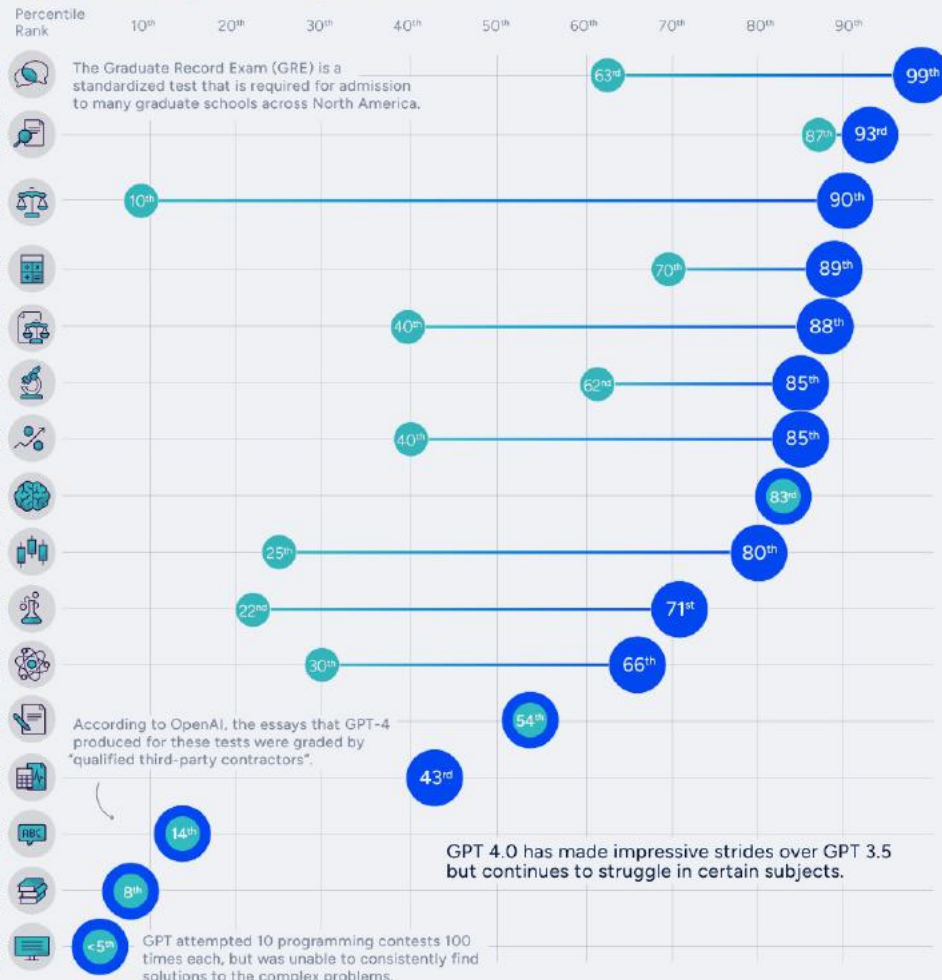


OpenAI's latest large language model, GPT-4, is capable of human-level performance in many professional and academic exams.

For example

60th Percentile

Exam Results ● ChatGPT 3.5 ● ChatGPT 4.0



GPT 4.0 has made impressive strides over GPT 3.5 but continues to struggle in certain subjects.

GPT attempted 10 programming contests 100 times each, but was unable to consistently find solutions to the complex problems.

Source: OpenAI (2023)

Note: Percentiles are based on the most recently available score distributions for test takers of each exam type.

  /visualcapitalist   @visualcap  visualcapitalist.com

14 March 2023

Question Answering

The first full-scale working railway steam locomotive was built by Richard Trevithick in the United Kingdom and, on 21 February 1804, the world's first railway journey took place as Trevithick's unnamed steam locomotive hauled a train along the tramway from the Pen-y-darren ironworks, near Merthyr Tydfil to Abercynon in south Wales. The design incorporated a number of important innovations that included using high-pressure steam which reduced the weight of the engine and increased its efficiency. Trevithick visited the Newcastle area later in 1804 and the colliery railways in north-east England became the leading centre for experimentation and development of steam locomotives.

In what country was a full-scale working railway steam locomotive first invented?

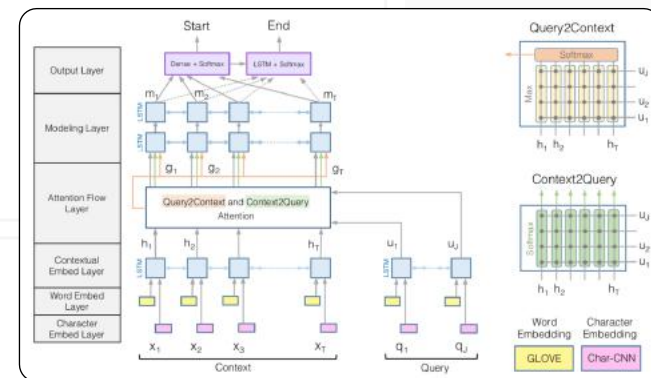
Ground Truth Answers: United Kingdom United Kingdom United Kingdom

Prediction: United Kingdom

On what date did the first railway trip in the world occur?

Ground Truth Answers: 21 February 1804 21 February 1804 21 February 1804

Prediction: 21 February 1804



Visual Question Answering



COCOQA 33827

What is the color of the cat?

Ground truth: black

IMG+BOW: black (0.55)

2-VIS+LSTM: black (0.73)

BOW: gray (0.40)

COCOQA 33827a

What is the color of the couch?

Ground truth: red

IMG+BOW: red (0.65)

2-VIS+LSTM: black (0.44)

BOW: red (0.39)



DAQUAR 1522

How many chairs are there?

Ground truth: two

IMG+BOW: four (0.24)

2-VIS+BLSTM: one (0.29)

LSTM: four (0.19)

DAQUAR 1520

How many shelves are there?

Ground truth: three

IMG+BOW: three (0.25)

2-VIS+BLSTM: two (0.48)

LSTM: two (0.21)



COCOQA 14855

Where are the ripe bananas sitting?

Ground truth: basket

IMG+BOW: basket (0.97)

2-VIS+BLSTM: basket (0.58)

BOW: bowl (0.48)

COCOQA 14855a

What are in the basket?

Ground truth: bananas

IMG+BOW: bananas (0.98)

2-VIS+BLSTM: bananas (0.68)

BOW: bananas (0.14)



DAQUAR 585

What is the object on the chair?

Ground truth: pillow

IMG+BOW: clothes (0.37)

2-VIS+BLSTM: pillow (0.65)

LSTM: clothes (0.40)

DAQUAR 585a

Where is the pillow found?

Ground truth: chair

IMG+BOW: bed (0.13)

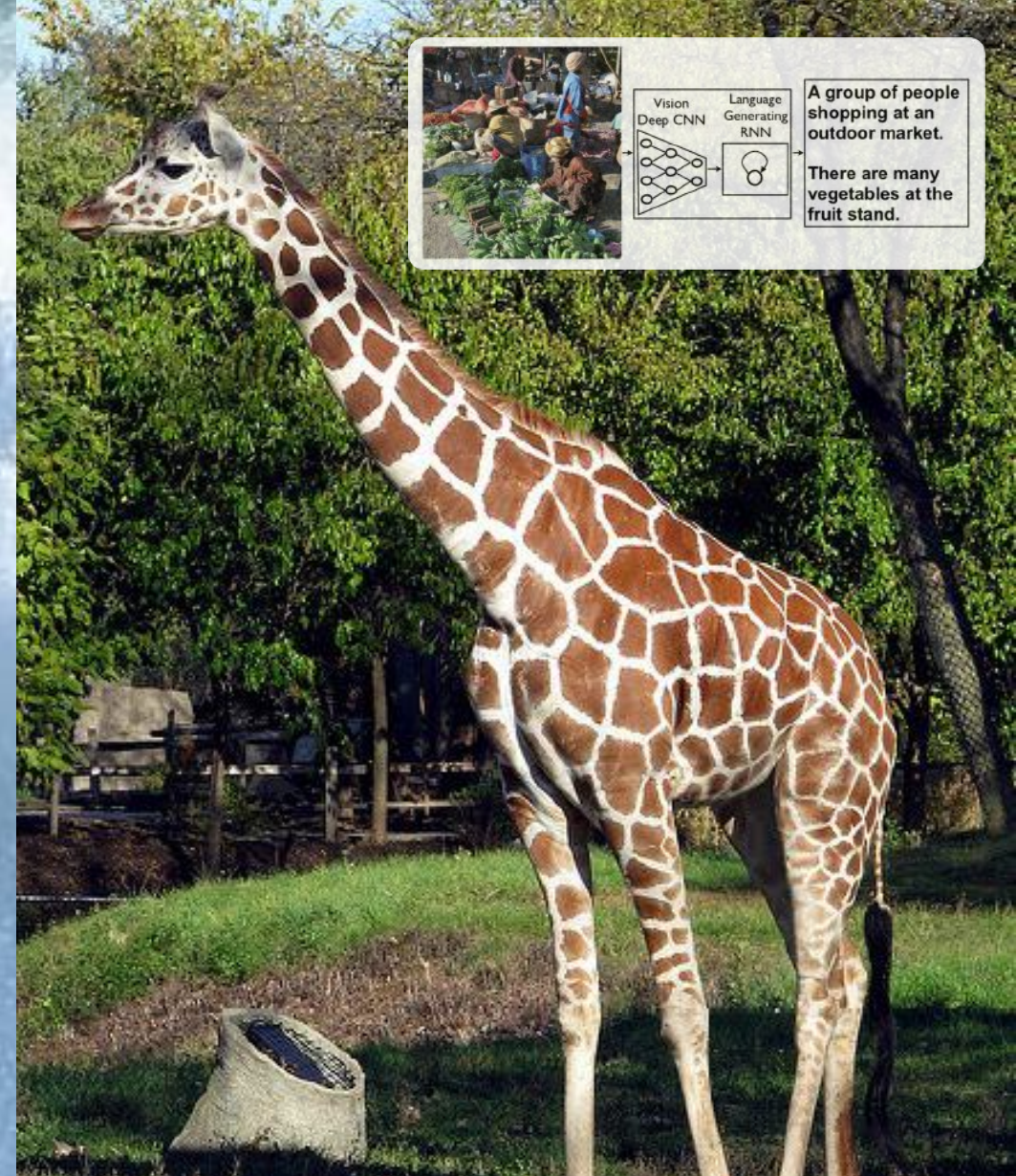
2-VIS+BLSTM: chair (0.17)

LSTM: cabinet (0.79)

Image Captioning



A man riding a wave on a surfboard in the water.



A giraffe standing in the grass next to a tree.

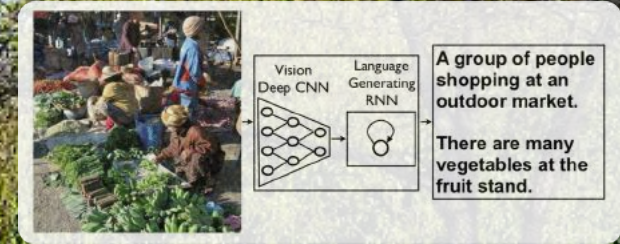
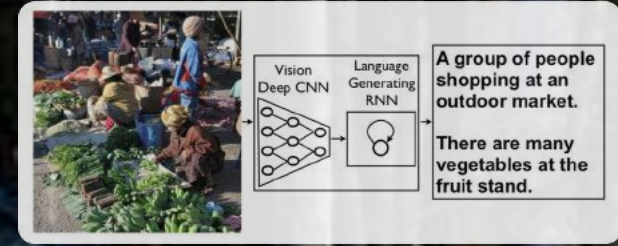


Image Captioning



Yarış pistinde virajı almakta olan bir yarış arabası

User What is unusual about this image?



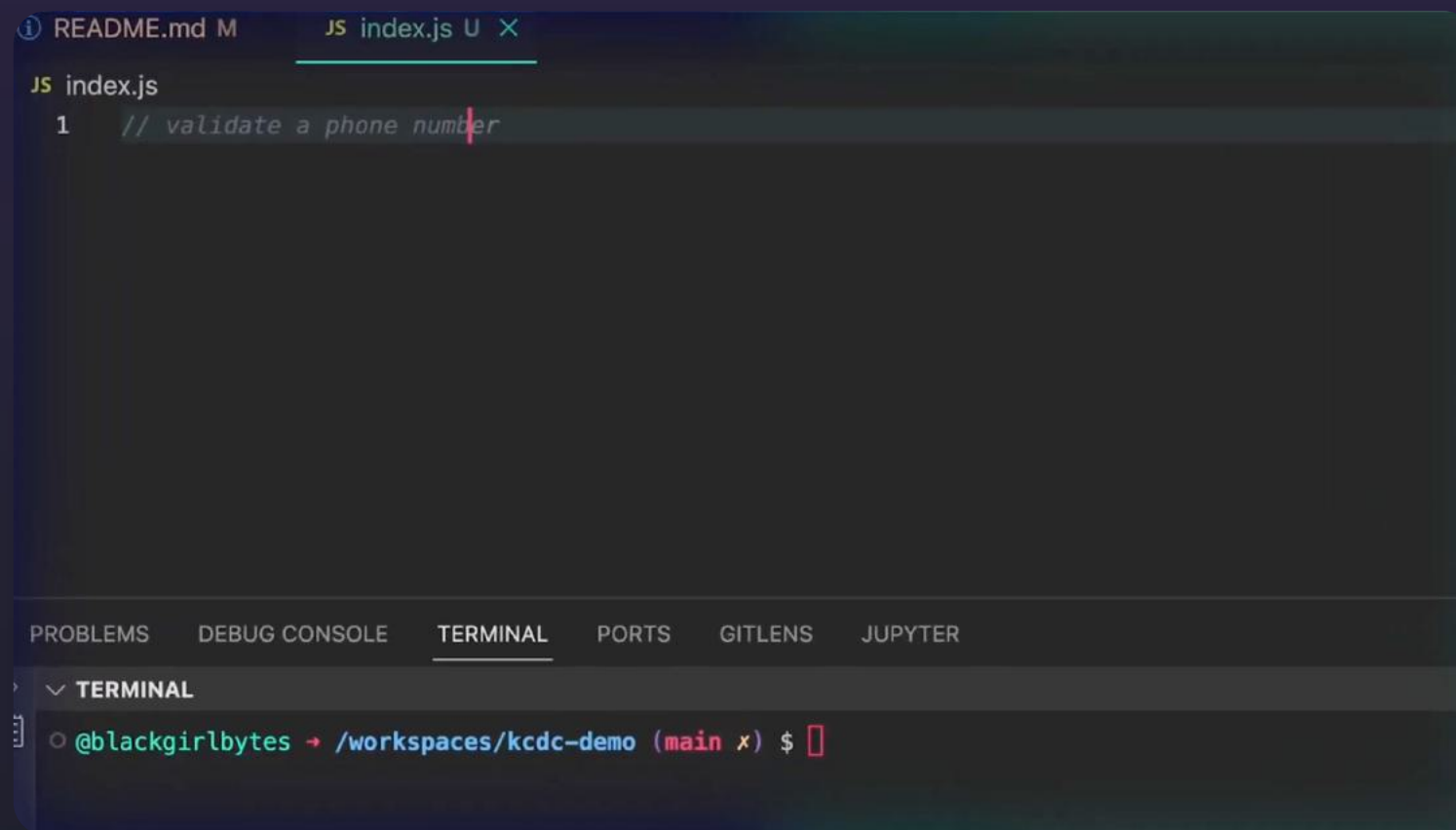
Source: [Barnorama](#)

GPT-4 The unusual thing about this image is that a man is ironing clothes on an ironing board attached to the roof of a moving taxi.



Your AI pair programmer

GitHub Copilot uses the OpenAI Codex to suggest code and entire functions in real-time, right from your editor.



Text Prompt

an armchair in the shape of an avocado. an armchair imitating an avocado.

AI generated
images



In the preceding visual, we explored DALL-E's ability to generate fantastical objects by combining two unrelated ideas. Here, we explore its ability to take inspiration from an unrelated idea while respecting the form of the thing being designed, ideally producing an object that appears to be practically functional. We found that prompting DALL-E with the phrases "in the shape of," "in the form of," and "in the style of" gives it the ability to do this.

When generating some of these objects, such as "an armchair in the shape of an avocado", DALL-E appears to relate the shape of a half avocado to the back of the chair, and the pit of the avocado to the cushion. We find that DALL-E is susceptible to the same kinds of mistakes mentioned in the previous visual.



A brain riding a rocketship heading towards the moon.



A photo of a Corgi dog riding a bike in Times Square. It is wearing sunglasses and a beach



A cute corgi lives in a house made out of sushi.



A blue jay standing on a large basket of rainbow macarons.



A transparent sculpture of a duck made out of glass.



A bald eagle made of chocolate powder, mango, and whipped cream.



An extremely angry bird.



A single beam of light enter the room from the ceiling. The beam of light is illuminating an easel. On the easel there is a Rembrandt painting of a raccoon.



A teddy bear
running in New York City



A british shorthair
jumping over a couch



A swarm of bees
flying around their hive



Melting pistachio ice cream
dripping down the cone.



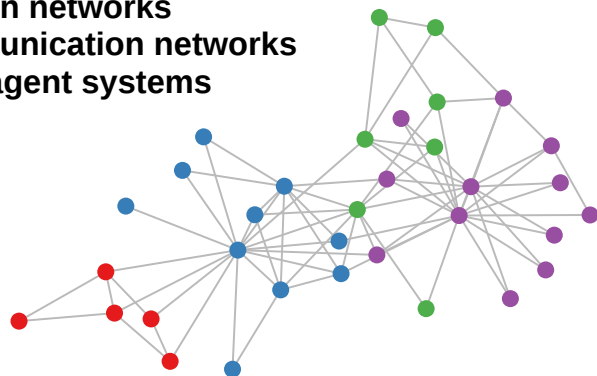
A british shorthair
jumping over a couch



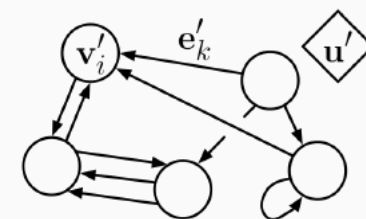
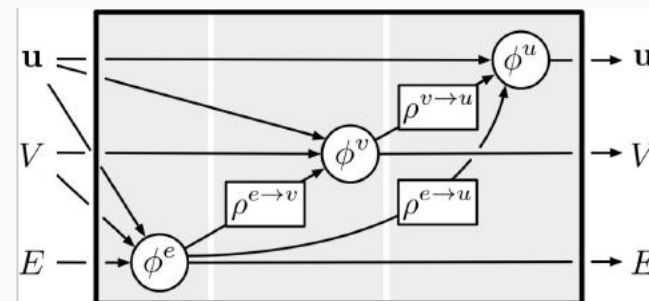
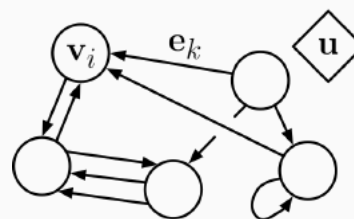
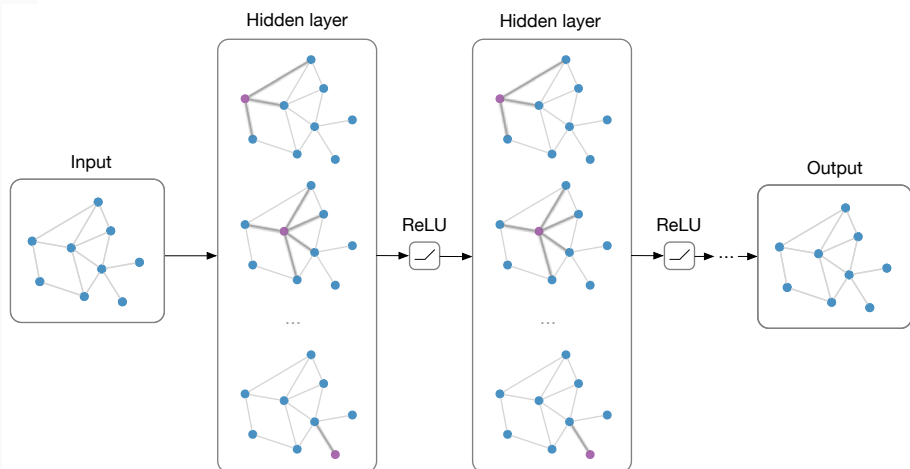
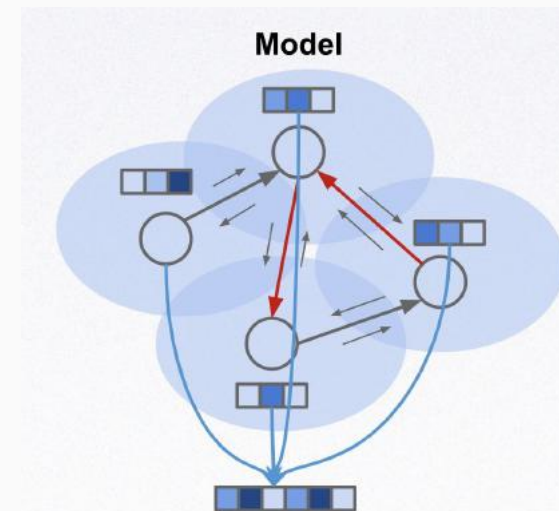
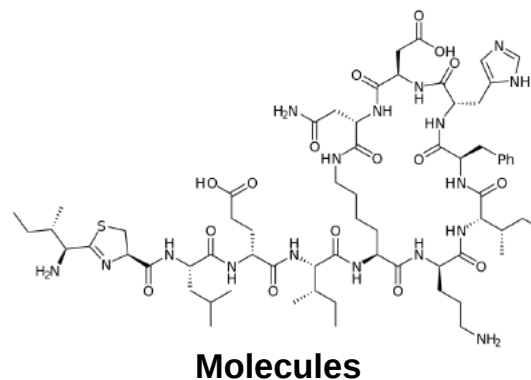
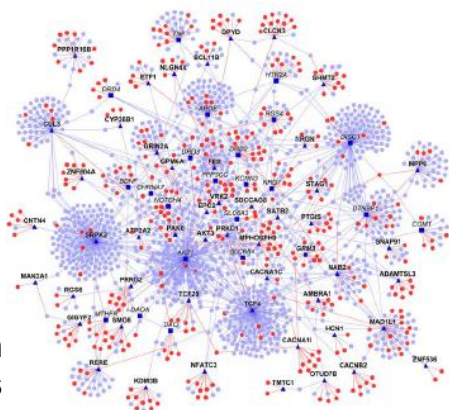
A shark swimming in clear
Carribean ocean.

Graph Neural Networks

Social networks
Citation networks
Communication networks
Multi-agent systems



Protein interaction networks



T.N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks", ICLR 2017

P. Battaglia et al., "Relational inductive biases, deep learning, and graph networks", arXiv 2018

linear-Gaussian controller training

welcomes

ROBOTICS
SCIENCE AND SYSTEMS

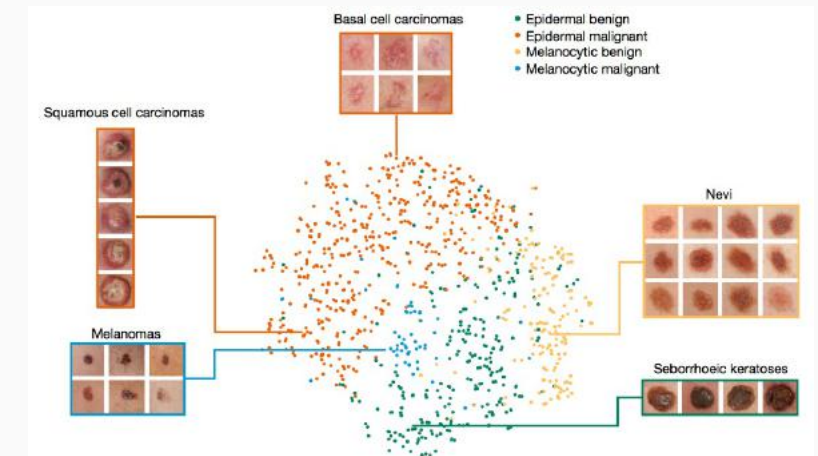
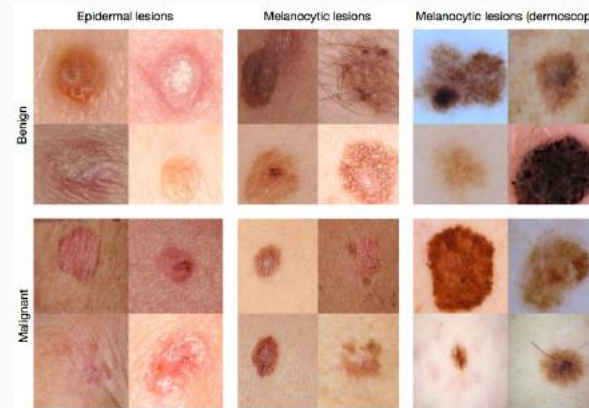
PR2

Robotics

autonomous execution

<http://rll.berkeley.edu/deeplearningrobotics/>

Medical Image Analysis

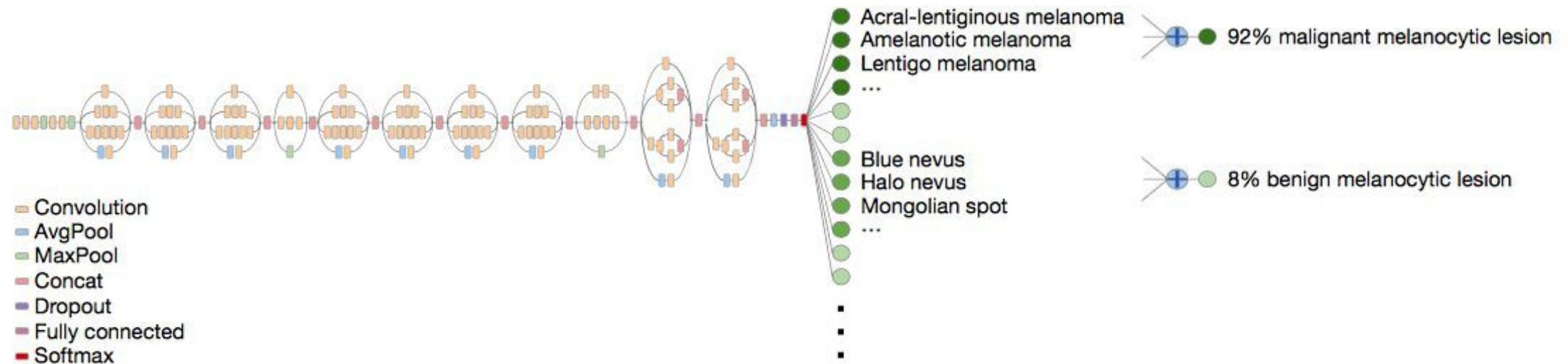


Skin lesion image

Deep convolutional neural network (Inception v3)

Training classes (757)

Inference classes (varies by task)



CheXNet: Radiologist-Level Pneumonia Detection on Chest X-Rays with Deep Learning

Pranav Rajpurkar*, Jeremy Irvin*, Kaylie Zhu, Brandon Yang, Hershel Mehta, Tony Duan, Daisy Ding, Aarti Bagul, Curtis Langlotz, Katie Shpanskaya, Matthew P. Lungren, Andrew Y. Ng

We develop an algorithm that can detect pneumonia from chest X-rays at a level exceeding practicing radiologists.

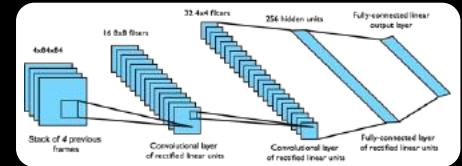
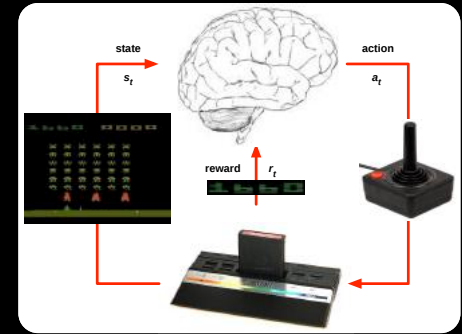
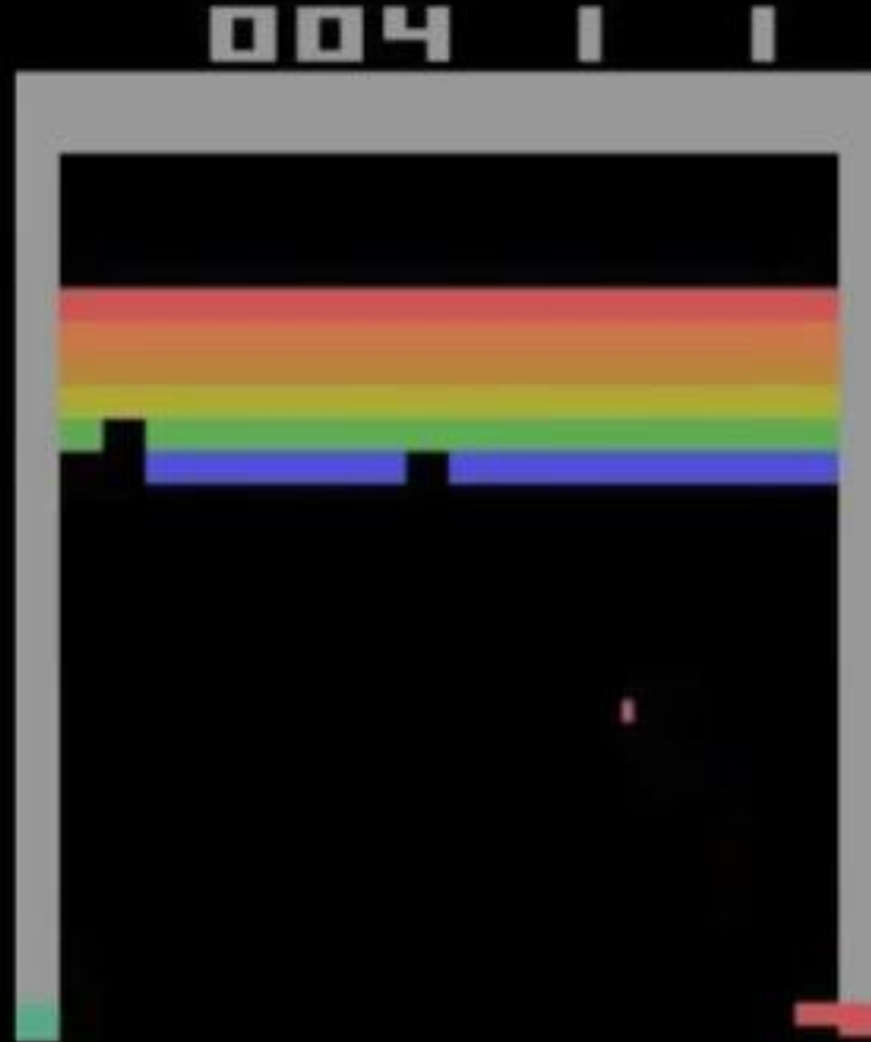
Chest X-rays are currently the best available method for diagnosing pneumonia, playing a crucial role in clinical care and epidemiological studies. Pneumonia is responsible for more than 1 million hospitalizations and 50,000 deaths per year in the US alone.

[READ OUR PAPER](#)

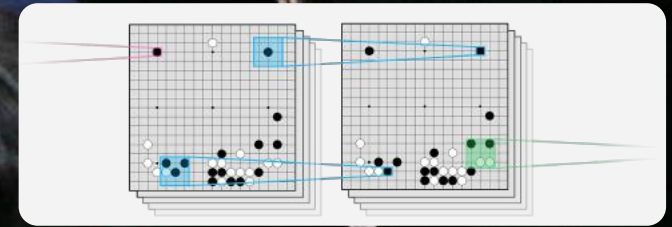
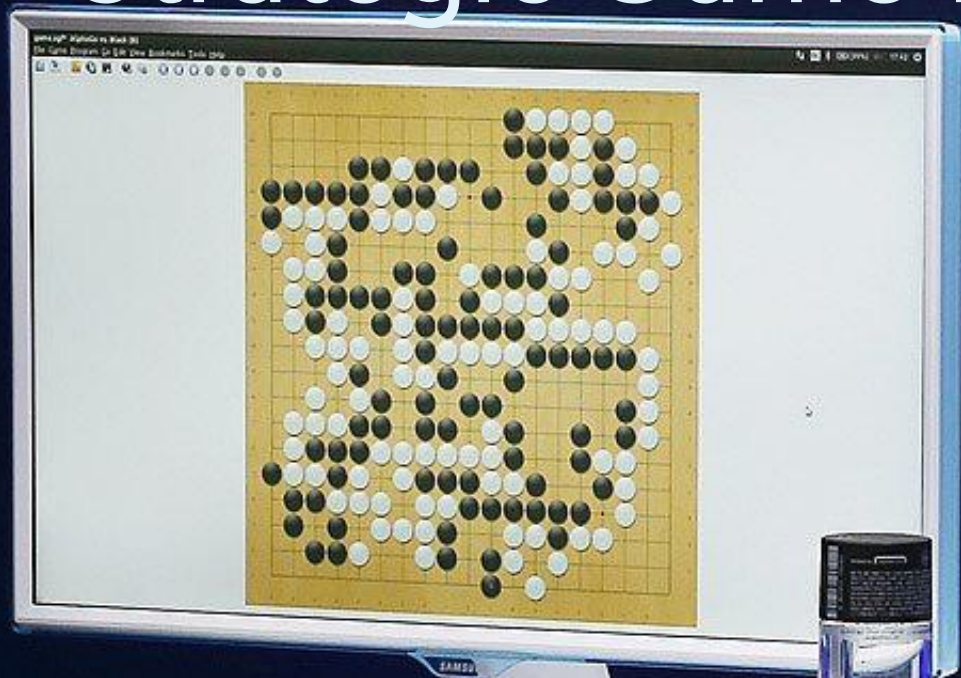


Medical Image Analysis

Strategic Game Playing



Strategic Game Playing



- AlphaGo vs. Lee Sidol
- Move 37, Game 2

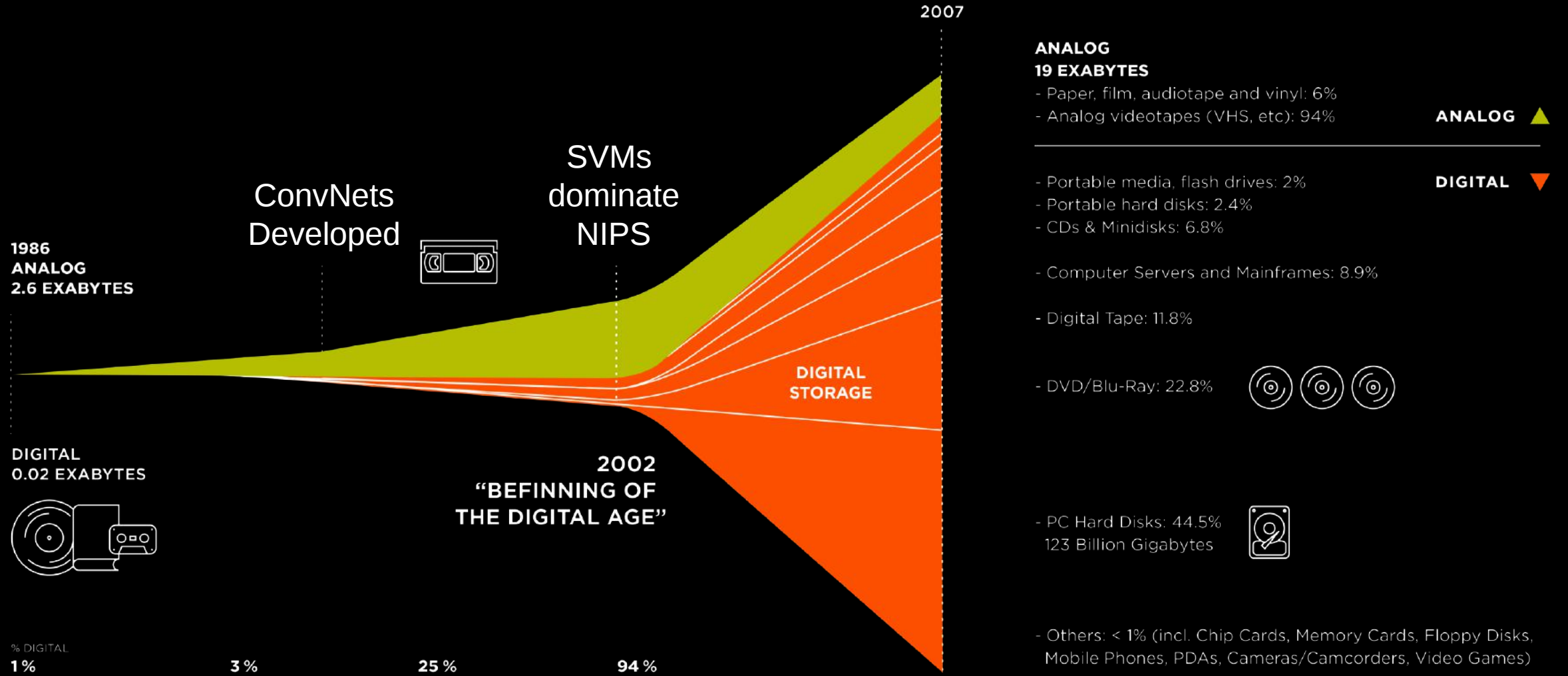
Bioinformatics



Why now?

The Resurgence of Deep Learning

GLOBAL INFORMATION STORAGE CAPACITY IN OPTIMALLY COMPRESSED BYTES



Source: Hilbert, M., & López, P. (2011). The World's Technological Capacity to Store, Communicate, and Compute Information. Science, 332 (6025), 60-65. martinhilbert.net/worldinfocapacity.html

Datasets vs. Algorithms

Year	Breakthroughs in AI	Datasets (First Available)	Algorithms (First Proposed)
1994	Human-level spontaneous speech recognition	Spoken Wall Street Journal articles and other texts (1991)	Hidden Markov Model (1984)
1997	IBM Deep Blue defeated Garry Kasparov	700,000 Grandmaster chess games, aka "The Extended Book" (1991)	Negascout planning algorithm (1983)
2005	Google's Arabic-and Chinese-to-English translation	1.8 trillion tokens from Google Web and News pages (collected in 2005)	Statistical machine translation algorithm (1988)
2011	IBM Watson became the world Jeopardy! champion	8.6 million documents from Wikipedia, Wiktionary, and Project Gutenberg (updated in 2010)	Mixture-of-Experts (1991)
2014	Google's GoogLeNet object classification at near-human performance	ImageNet corpus of 1.5 million labeled images and 1,000 object categories (2010)	Convolutional Neural Networks (1989)
2015	Google's DeepMind achieved human parity in playing 29 Atari games by learning general control from video	Arcade Learning Environment dataset of over 50 Atari games (2013)	Q-learning (1992)
Average No. of Years to Breakthrough:		3 years	18 years

Powerful Hardware

- Deep neural nets highly amenable to implementation on Graphics Processing Units (GPUs)
 - Matrix multiplication
 - 2D convolution
- E.g. nVidia Pascal GPUs deliver 10 Tflops
 - Faster than fastest computer in the world in 2000
 - 10 million times faster than 1980's Sun workstation

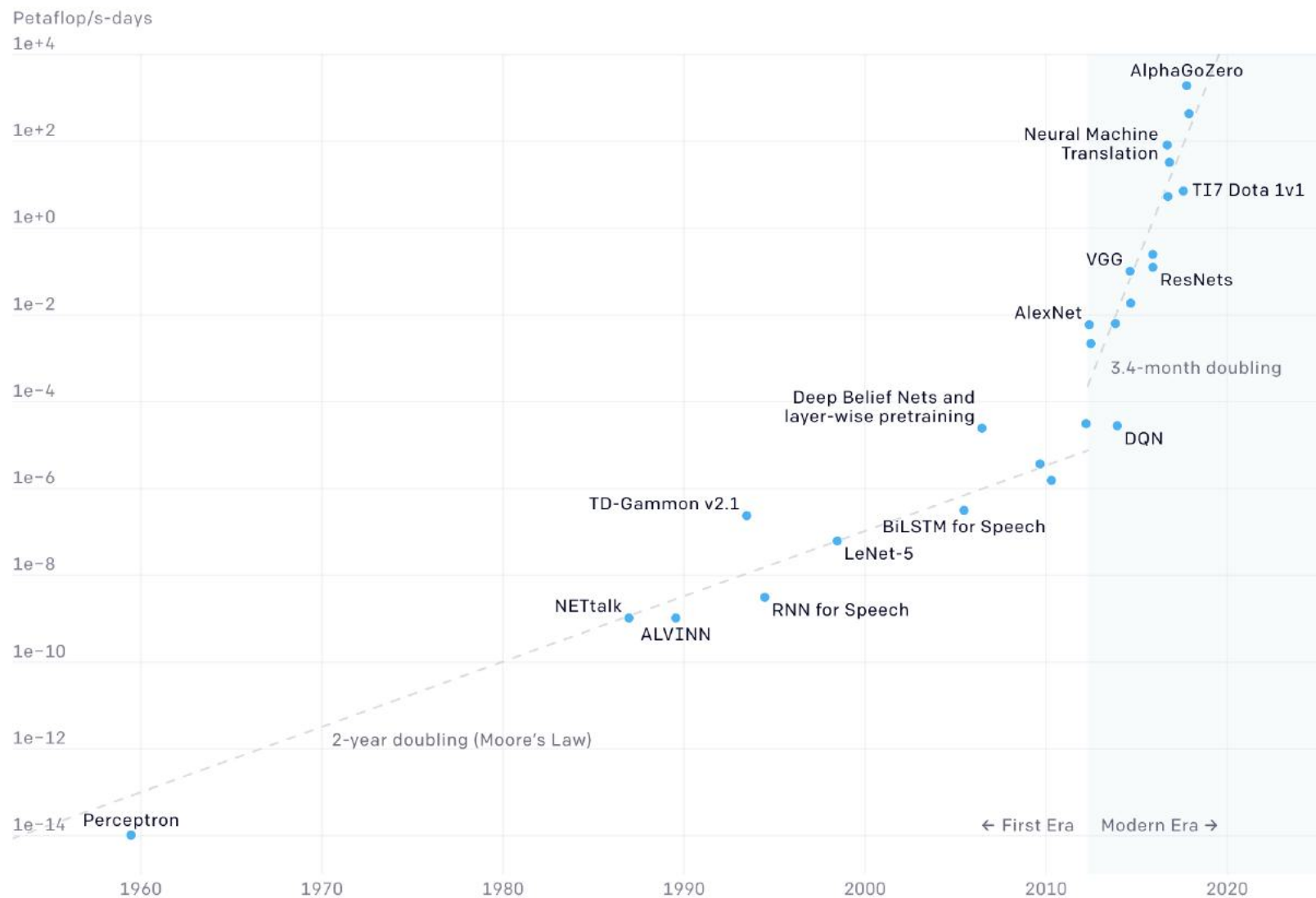


Image: OpenAI

Working ideas on how to train deep architectures

Dropout: A Simple Way to Prevent Neural Networks from Overfitting

Nitish Srivastava
Geoffrey Hinton
Alex Krizhevsky
Ilya Sutskever
Ruslan Salakhutdinov

NITISH@CS.TORONTO.EDU
HINTON@CS.TORONTO.EDU
KRIZ@CS.TORONTO.EDU
ILYA@CS.TORONTO.EDU
RSALAKHU@CS.TORONTO.EDU

Abstract

Deep neural nets with a large number of parameters are very powerful machine learning systems. However, overfitting is a serious problem in such networks. Large networks are also slow to use, making it difficult to deal with overfitting by combining the predictions of many different large neural nets at test time. Dropout is a technique for addressing this problem. The key idea is to randomly drop units (along with their connections) from the neural network during training. This prevents units from co-adapting too much. During training, dropout samples from an exponential number of different “thinned” networks. At test time,

Journal of Machine Learning Research 15 (2014) 1929-1958

Submitted 11/13, Published 6/14

Dropout: A Simple Way to Prevent Neural Networks from Overfitting

Nitish Srivastava
Geoffrey Hinton
Alex Krizhevsky
Ilya Sutskever
Ruslan Salakhutdinov
*Department of Computer Science
University of Toronto
10 King's College Road, Rm 3308
Toronto, Ontario, M5S 3G4, Canada.*

NITISH@CS.TORONTO.EDU
HINTON@CS.TORONTO.EDU
KRIZ@CS.TORONTO.EDU
ILYA@CS.TORONTO.EDU
RSALAKHU@CS.TORONTO.EDU

Editor: Yoshua Bengio

Abstract

Deep neural nets with a large number of parameters are very powerful machine learning systems. However, overfitting is a serious problem in such networks. Large networks are also slow to use, making it difficult to deal with overfitting by combining the predictions of many different large neural nets at test time. Dropout is a technique for addressing this problem. The key idea is to randomly drop units (along with their connections) from the neural network during training. This prevents units from co-adapting too much. During training, dropout samples from an exponential number of different “thinned” networks. At test time, it is easy to approximate the effect of averaging the predictions of all these thinned networks by simply using a single unthinned network that has smaller weights. This significantly reduces overfitting and gives major improvements over other regularization methods. We show that dropout improves the performance of neural networks on supervised learning tasks in vision, speech recognition, document classification and computational biology, obtaining state-of-the-art results on many benchmark data sets.

Keywords: neural networks, regularization, model combination, deep learning

1. Introduction

Deep neural networks contain multiple non-linear hidden layers and this makes them very expressive models that can learn very complicated relationships between their inputs and outputs. With limited training data, however, many of these complicated relationships will be the result of sampling noise, so they will exist in the training set but not in real test data even if it is drawn from the same distribution. This leads to overfitting and many methods have been developed for reducing it. These include stopping the training as soon as performance on a validation set starts to get worse, introducing weight penalties of various kinds such as L1 and L2 regularization and soft weight sharing (Nowlan and Hinton, 1992).

With unlimited computation, the best way to “regularize” a fixed-sized model is to average the predictions of all possible settings of the parameters, weighting each setting by

©2014 Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever and Ruslan Salakhutdinov.

- Better Learning Regularization (e.g. **Dropout**)

N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, R. Salakhutdinov, “Dropout: A Simple Way to Prevent Neural Networks from Overfitting”, JMLR Vol. 15, No. 1,

Working ideas on how to train deep architectures

Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift

Sergey Ioffe
Google Inc., sioffe@google.com

Christian Szegedy
Google Inc., szegedy@google.com

Abstract

Training Deep Neural Networks is complicated by the fact that the distribution of each layer's inputs changes during training, as the parameters of the previous layers change. This slows down the training by requiring lower learning rates and careful parameter initialization, and makes it notoriously hard to train models with saturating nonlinearities. We refer to this phenomenon as *internal covariate shift*, and address the problem by normalizing layer inputs. Our method draws its strength from making normalization a part of the model architecture and performing the normalization for each training mini-batch. Batch Normalization

Using mini-batches of examples, as opposed to one example at a time, is helpful in several ways. First, the gradient of the loss over a mini-batch is an estimate of the gradient over the training set, whose quality improves as the batch size increases. Second, computation over a batch can be much more efficient than m computations for individual examples, due to the parallelism afforded by the modern computing platforms.

While stochastic gradient is simple and effective, it requires careful tuning of the model hyper-parameters, specifically the learning rate used in optimization, as well as the initial values for the model parameters. The training is complicated by the fact that the inputs to each layer

Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift

Sergey Ioffe
Google Inc., sioffe@google.com

Christian Szegedy
Google Inc., szegedy@google.com

Abstract

Training Deep Neural Networks is complicated by the fact that the distribution of each layer's inputs changes during training, as the parameters of the previous layers change. This slows down the training by requiring lower learning rates and careful parameter initialization, and makes it notoriously hard to train models with saturating nonlinearities. We refer to this phenomenon as *internal covariate shift*, and address the problem by normalizing layer inputs. Our method draws its strength from making normalization a part of the model architecture and performing the normalization for each training mini-batch. Batch Normalization allows us to use much higher learning rates and be less careful about initialization. It also acts as a regularizer, in some cases eliminating the need for Dropout. Applied to a state-of-the-art image classification model, Batch Normalization achieves the same accuracy with 14 times fewer training steps, and beats the original model by a significant margin. Using an ensemble of batch-normalized networks, we improve upon the best published result on ImageNet classification: reaching 4.9% top-5 validation error (and 4.8% test error), exceeding the accuracy of human raters.

1 Introduction

Deep learning has dramatically advanced the state of the art in vision, speech, and many other areas. Stochastic gradient descent (SGD) has proved to be an effective way of training deep networks, and SGD variants such as momentum (Sutskever et al., 2013) and Adagrad (Duchi et al., 2011) have been used to achieve state of the art performance. SGD optimizes the parameters Θ of the network, so as to minimize the loss

$$\Theta = \arg \min_{\Theta} \frac{1}{N} \sum_{i=1}^N \ell(x_i, \Theta)$$

where $x_{1..N}$ is the training data set. With SGD, the training proceeds in steps, and at each step we consider a mini-batch $x_{1..m}$ of size m . The mini-batch is used to approximate the gradient of the loss function with respect to the parameters, by computing

$$\frac{1}{m} \frac{\partial \ell(x_i, \Theta)}{\partial \Theta}$$

Using mini-batches of examples, as opposed to one example at a time, is helpful in several ways. First, the gradient of the loss over a mini-batch is an estimate of the gradient over the training set, whose quality improves as the batch size increases. Second, computation over a batch can be much more efficient than m computations for individual examples, due to the parallelism afforded by the modern computing platforms.

While stochastic gradient is simple and effective, it requires careful tuning of the model hyper-parameters, specifically the learning rate used in optimization, as well as the initial values for the model parameters. The training is complicated by the fact that the inputs to each layer are affected by the parameters of all preceding layers – so that small changes to the network parameters amplify as the network becomes deeper.

The change in the distributions of layers' inputs presents a problem because the layers need to continuously adapt to the new distribution. When the input distribution to a learning system changes, it is said to experience *covariate shift* (Shimodaira, 2000). This is typically handled via domain adaptation (Jiang, 2008). However, the notion of covariate shift can be extended beyond the learning system as a whole, to apply to its parts, such as a sub-network or a layer. Consider a network computing

$$\ell = F_2(F_1(x, \Theta_1), \Theta_2)$$

where F_1 and F_2 are arbitrary transformations, and the parameters Θ_1, Θ_2 are to be learned so as to minimize the loss ℓ . Learning Θ_2 can be viewed as if the inputs $x = F_1(x, \Theta_1)$ are fed into the sub-network

$$\ell = F_2(x, \Theta_2).$$

For example, a gradient descent step

$$\Theta_2 \leftarrow \Theta_2 - \alpha \sum_{i=1}^m \frac{\partial F_2(x_i, \Theta_2)}{\partial \Theta_2}$$

(for batch size m and learning rate α) is exactly equivalent to that for a stand-alone network F_2 with input x . Therefore, the input distribution properties that make training more efficient – such as having the same distribution between the training and test data – apply to training the sub-network as well. As such it is advantageous for the distribution of x to remain fixed over time. Then, Θ_2 does

- Better Optimization Conditioning (e.g. **Batch Normalization**)

Working ideas on how to train deep architectures

Deep Residual Learning for Image Recognition

Kaiming He Xiangyu Zhang Shaoqing Ren Jian Sun

Microsoft Research

{kahe, v-xiangz, v-shren, jiansun}@microsoft.com

Abstract

Deeper neural networks are more difficult to train. We present a residual learning framework to ease the training of networks that are substantially deeper than those used previously. We explicitly reformulate the layers as learning residual functions with reference to the layer inputs, instead of learning unreferenced functions. We provide comprehensive empirical evidence showing that these residual networks are easier to optimize, and can gain accuracy from considerably increased depth. On the ImageNet dataset we evaluate residual nets with a depth of up to 152 layers—8× deeper than VGG nets [41] but still having lower complexity. An ensemble of these residual nets achieves 3.57% error

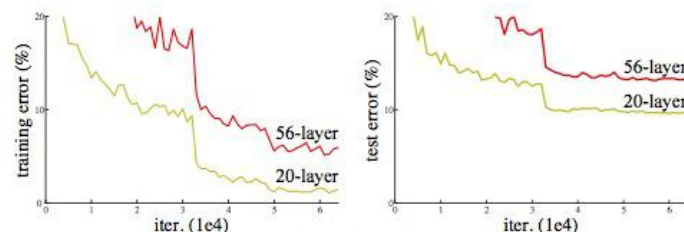


Figure 1. Training error (left) and test error (right) on CIFAR-10 with 20-layer and 56-layer “plain” networks. The deeper network has higher training error, and thus test error. Similar phenomena on ImageNet is presented in Fig. 4.

greatly benefited from very deep models.

Driven by the significance of depth, a question arises: Is

Deep Residual Learning for Image Recognition

Kaiming He Xiangyu Zhang Shaoqing Ren Jian Sun
Microsoft Research
{kahe, v-xiangz, v-shren, jiansun}@microsoft.com

Abstract

Deeper neural networks are more difficult to train. We present a residual learning framework to ease the training of networks that are substantially deeper than those used previously. We explicitly reformulate the layers as learning residual functions with reference to the layer inputs, instead of learning unreferenced functions. We provide comprehensive empirical evidence showing that these residual networks are easier to optimize, and can gain accuracy from considerably increased depth. On the ImageNet dataset we evaluate residual nets with a depth of up to 152 layers—8× deeper than VGG nets [41] but still having lower complexity. An ensemble of these residual nets achieves 3.57% error on the ImageNet test set. This result won the 1st place on the ILSVRC 2015 classification task. We also present analysis on CIFAR-10 with 100 and 1000 layers.

The depth of representations is of central importance for many visual recognition tasks. Solely due to our extremely deep representations, we obtain a 28% relative improvement on the COCO object detection dataset. Deep residual nets are foundations of our submissions to ILSVRC & COCO 2015 competitions¹, where we also won the 1st places on the tasks of ImageNet detection, ImageNet localization, COCO detection, and COCO segmentation.

1. Introduction

Deep convolutional neural networks [22, 21] have led to a series of breakthroughs for image classification [21, 50, 40]. Deep networks naturally integrate low/mid/high-level features [50] and classifiers in an end-to-end multi-layer fashion, and the “levels” of features can be enriched by the number of stacked layers (depth). Recent evidence [41, 44] reveals that network depth is of crucial importance, and the leading results [41, 44, 13, 16] on the challenging ImageNet dataset [36] all exploit “very deep” [41] models, with a depth of sixteen [41] to thirty [16]. Many other non-trivial visual recognition tasks [8, 12, 7, 32, 27] have also

¹<http://image-net.org/challenges/LSVRC/2015/> and <http://mscoco.org/dataset/#detection-challenge2015>.

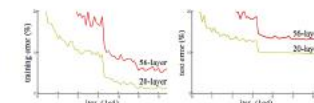


Figure 1. Training error (left) and test error (right) on CIFAR-10 with 20-layer and 56-layer “plain” networks. The deeper network has higher training error, and thus test error. Similar phenomena on ImageNet is presented in Fig. 4.

greatly benefited from very deep models.

Driven by the significance of depth, a question arises: Is learning better networks as easy as stacking more layers? An obstacle to answering this question was the notorious problem of vanishing/exploding gradients [1, 9], which hamper convergence from the beginning. This problem, however, has been largely addressed by normalized initialization [23, 9, 37, 13] and intermediate normalization layers [16], which enable networks with tens of layers to start converging for stochastic gradient descent (SGD) with back-propagation [22].

When deeper networks are able to start converging, a degradation problem has been exposed: with the network depth increasing, accuracy gets saturated (which might be unsurprising) and then degrades rapidly. Unexpectedly, such degradation is not caused by overfitting, and adding more layers to a suitably deep model leads to higher training error, as reported in [11, 42] and thoroughly verified by our experiments. Fig. 1 shows a typical example.

The degradation (of training accuracy) indicates that not all systems are similarly easy to optimize. Let us consider a shallower architecture and its deeper counterpart that adds more layers onto it. There exists a solution by construction to the deeper model: the added layers are identity mapping, and the other layers are copied from the learned shallower model. The existence of this constructed solution indicates that a deeper model should produce no higher training error than its shallower counterpart. But experiments show that our current solvers on hand are unable to find solutions that

- Better neural architectures (e.g. **Residual Nets**)

Software



Caffe



Caffe2

MatConvNet

The Microsoft Cognitive Toolkit

A free, easy-to-use, open-source, commercial-grade toolkit that trains deep learning algorithms to learn like the human brain.



PYTORCH

So what is deep learning?

Three key ideas

- (Hierarchical) Compositionality
 - Cascade of non-linear transformations
 - Multiple layers of representations
- End-to-End Learning
 - Learning (goal-driven) representations
 - Learning to feature extract
- Distributed Representations
 - No single neuron “encodes” everything
 - Groups of neurons work together

Three key ideas

- **(Hierarchical) Compositionality**
 - Cascade of non-linear transformations
 - Multiple layers of representations
- End-to-End Learning
 - Learning (goal-driven) representations
 - Learning to feature extract
- Distributed Representations
 - No single neuron “encodes” everything
 - Groups of neurons work together

Traditional Machine Learning

VISION



hand-crafted
features
SIFT/HOG

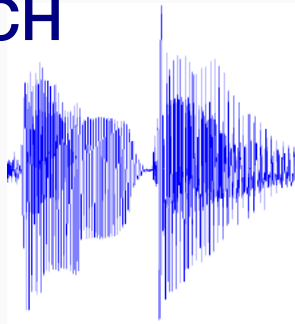
fixed

your favorite
classifier

learned

“car”

SPEECH



hand-crafted
features
MFCC

fixed

your favorite
classifier

learned

\ 'd ē p \

NLP

This burrito place
is yummy and fun!

hand-crafted
features
Bag-of-words

fixed

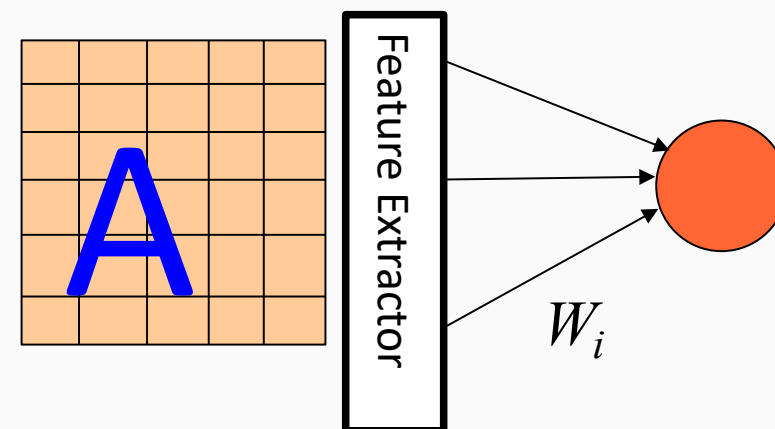
your favorite
classifier

learned

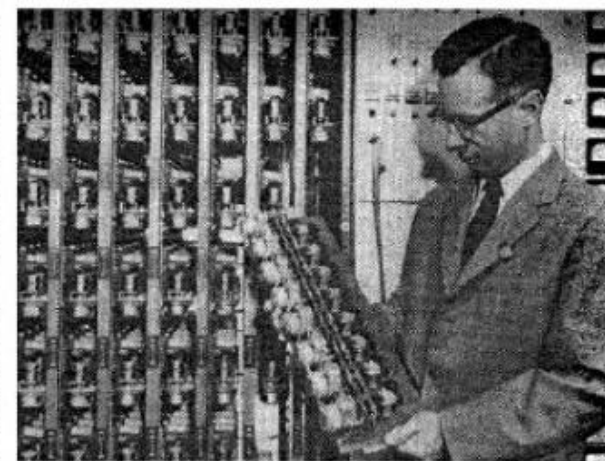
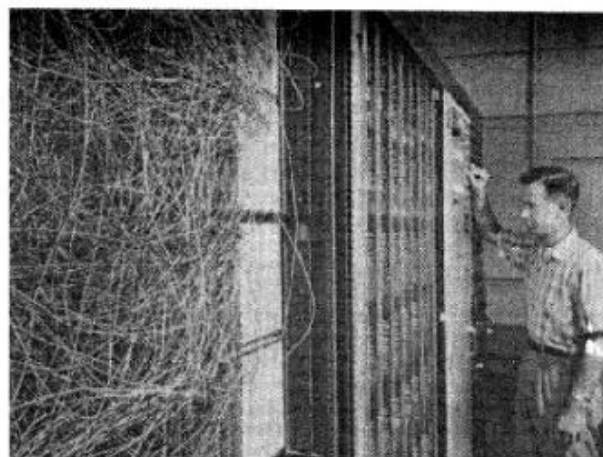
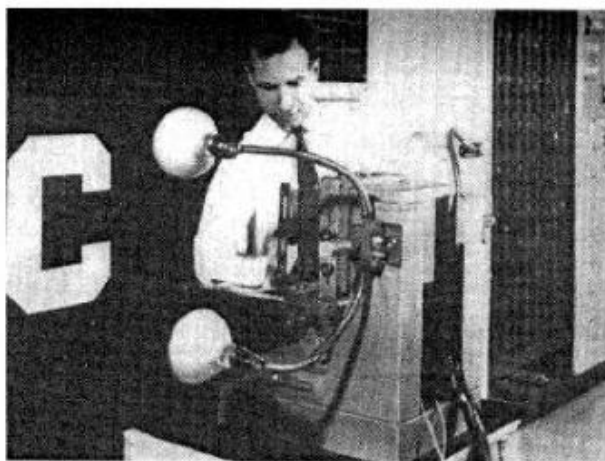
“+”

It's an old paradigm

- The first learning machine: the **Perceptron**
 - Built at Cornell in 1960
- The Perceptron was a **linear classifier** on top of a simple **feature extractor**
- The vast majority of practical applications of ML today use glorified **linear classifiers** or glorified template matching.
- Designing a feature extractor requires considerable efforts by experts.



$$y = \text{sign} \left(\sum_i^N W_i F_i(X) + b \right)$$



Hierarchical Compositionality

VISION

pixels → edge → texton → motif → part → object

SPEECH

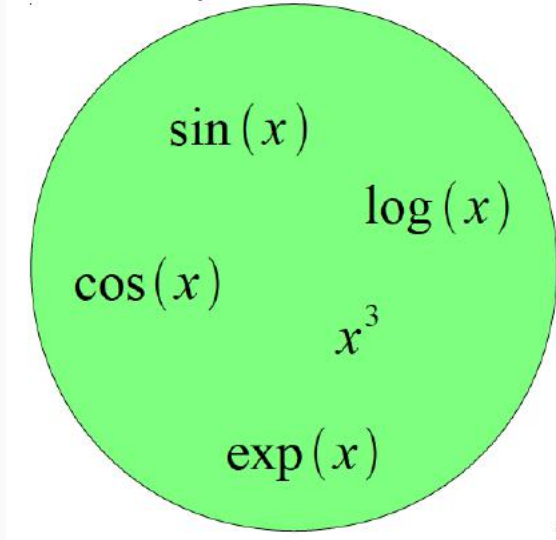
sample → spectral
band → formant → motif → phone → word

NLP

character → word → NP/VP/.. → clause → sentence → story

Building A Complicated Function

Given a library of simple functions

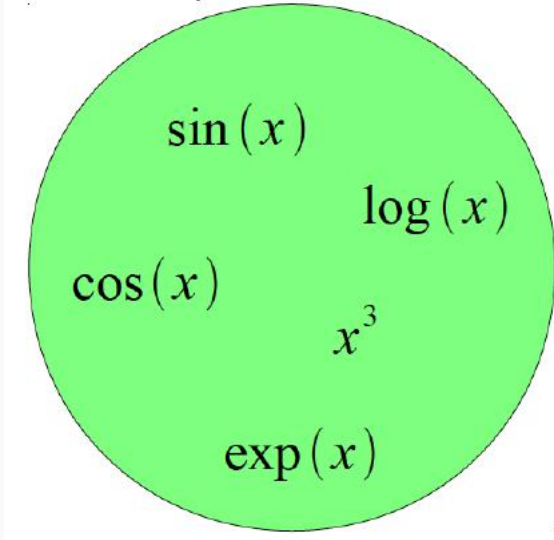


Compose into a

complicate function

Building A Complicated Function

Given a library of simple functions

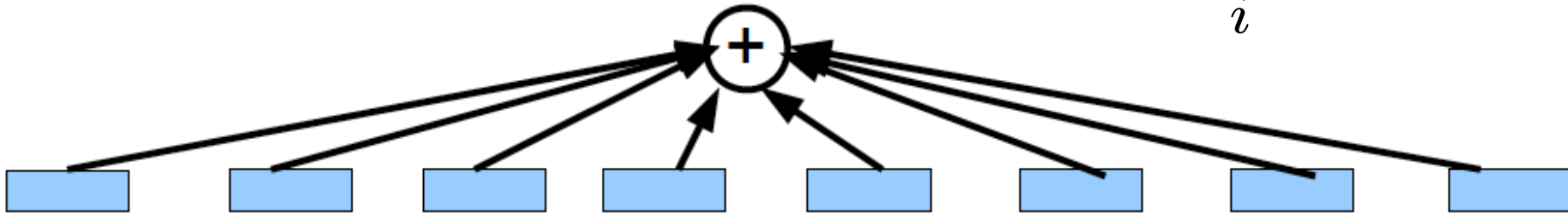


Compose into a
complicate function

Idea 1: Linear Combinations

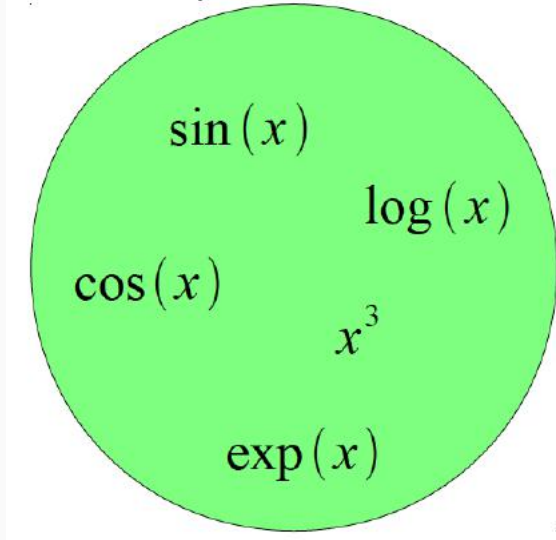
- Boosting
- Kernels
- ...

$$f(x) = \sum_i \alpha_i g_i(x)$$



Building A Complicated Function

Given a library of simple functions



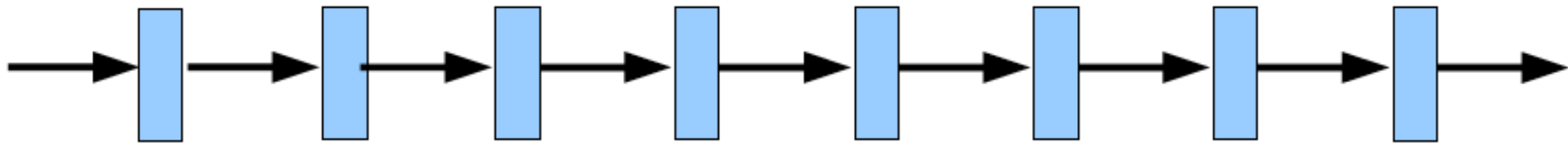
Compose into a

complicate function

Idea 2: Compositions

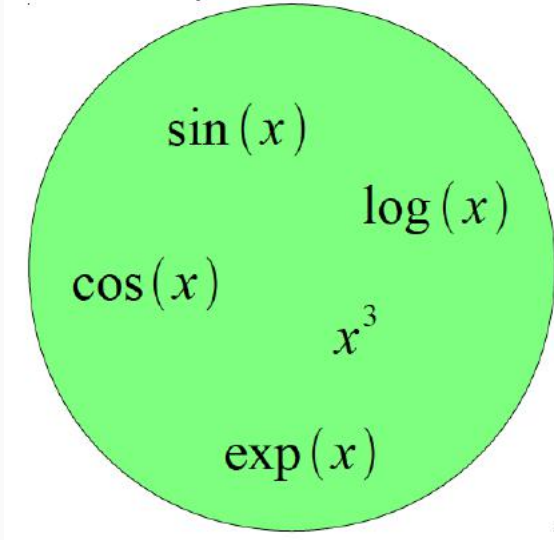
- Deep Learning
- Grammar models
- Scattering transforms...

$$f(x) = g_1(g_2(\dots(g_n(x)\dots)))$$



Building A Complicated Function

Given a library of simple functions

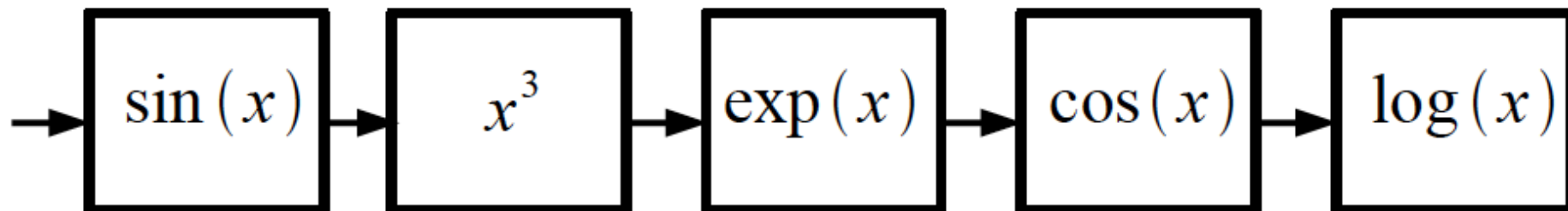


Compose into a
→
complicate function

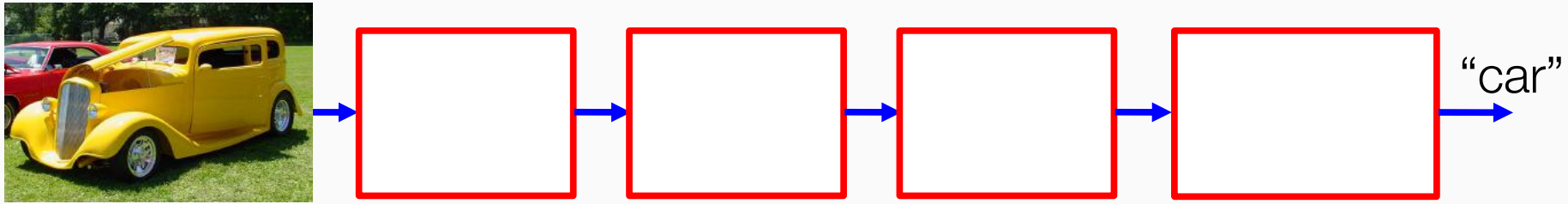
Idea 2: Compositions

- Deep Learning
- Grammar models
- Scattering transforms...

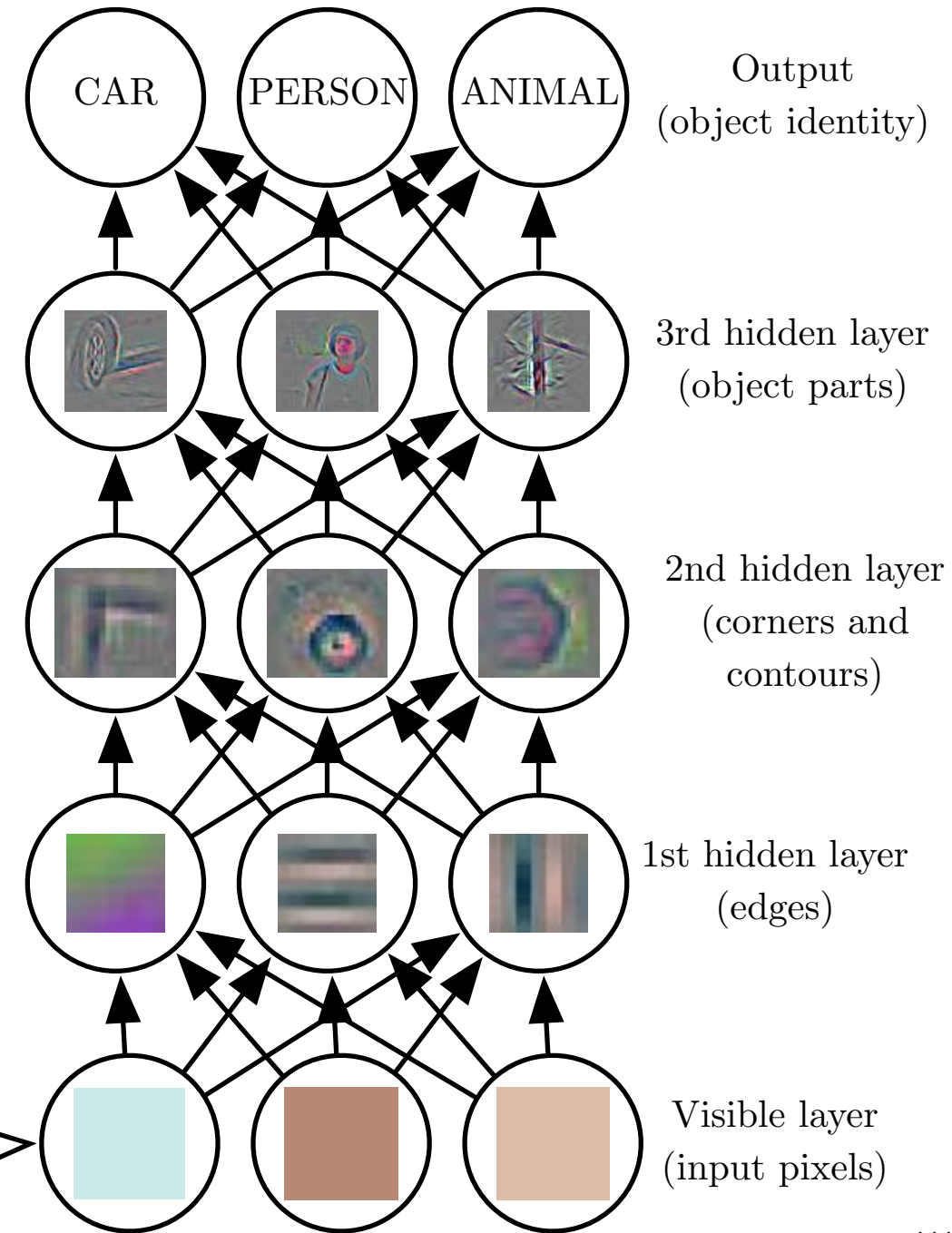
$$f(x) = \log(\cos(\exp(\sin^3(x))))$$



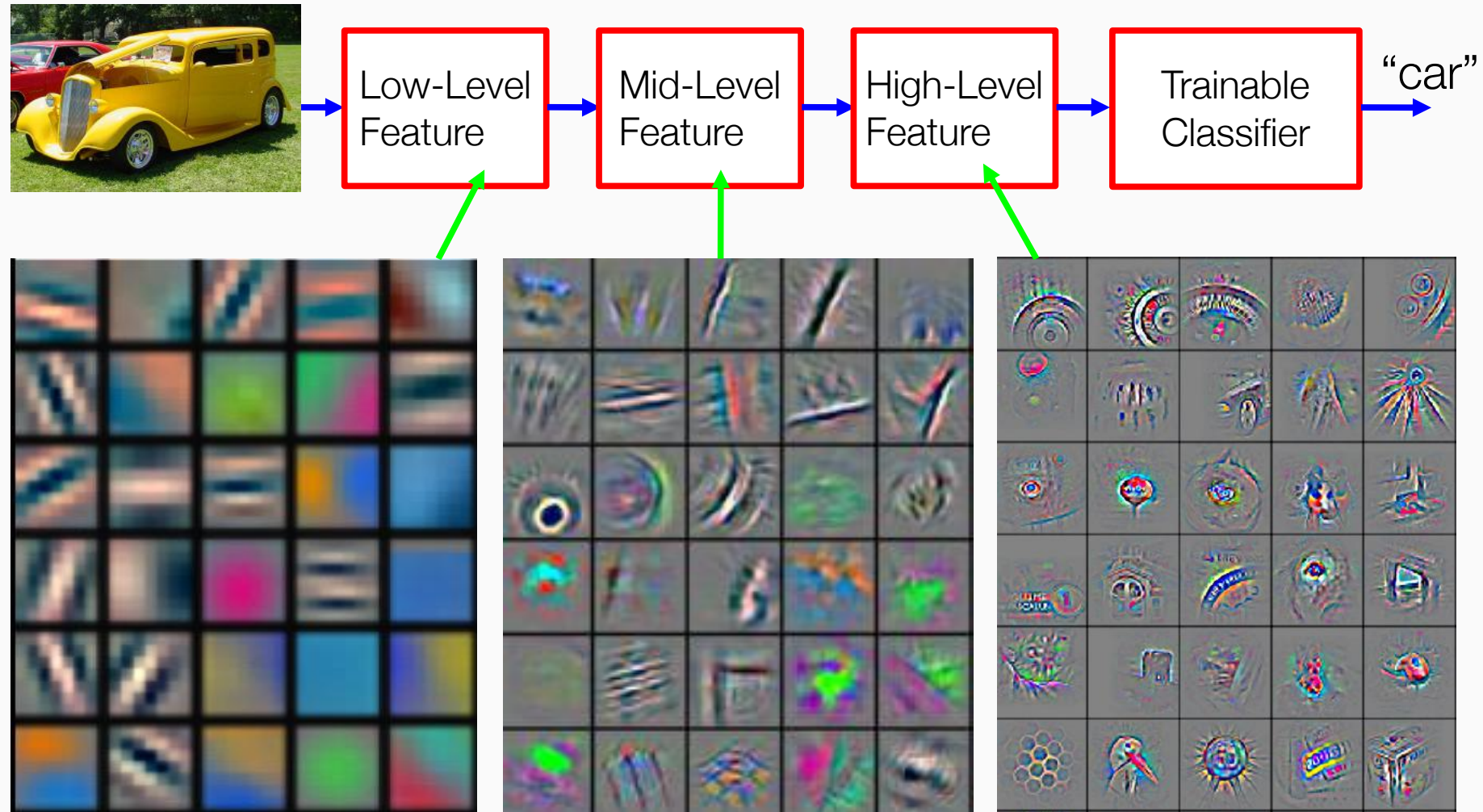
Deep Learning = Hierarchical Compositionality



Deep Learning = Hierarchical Compositionality

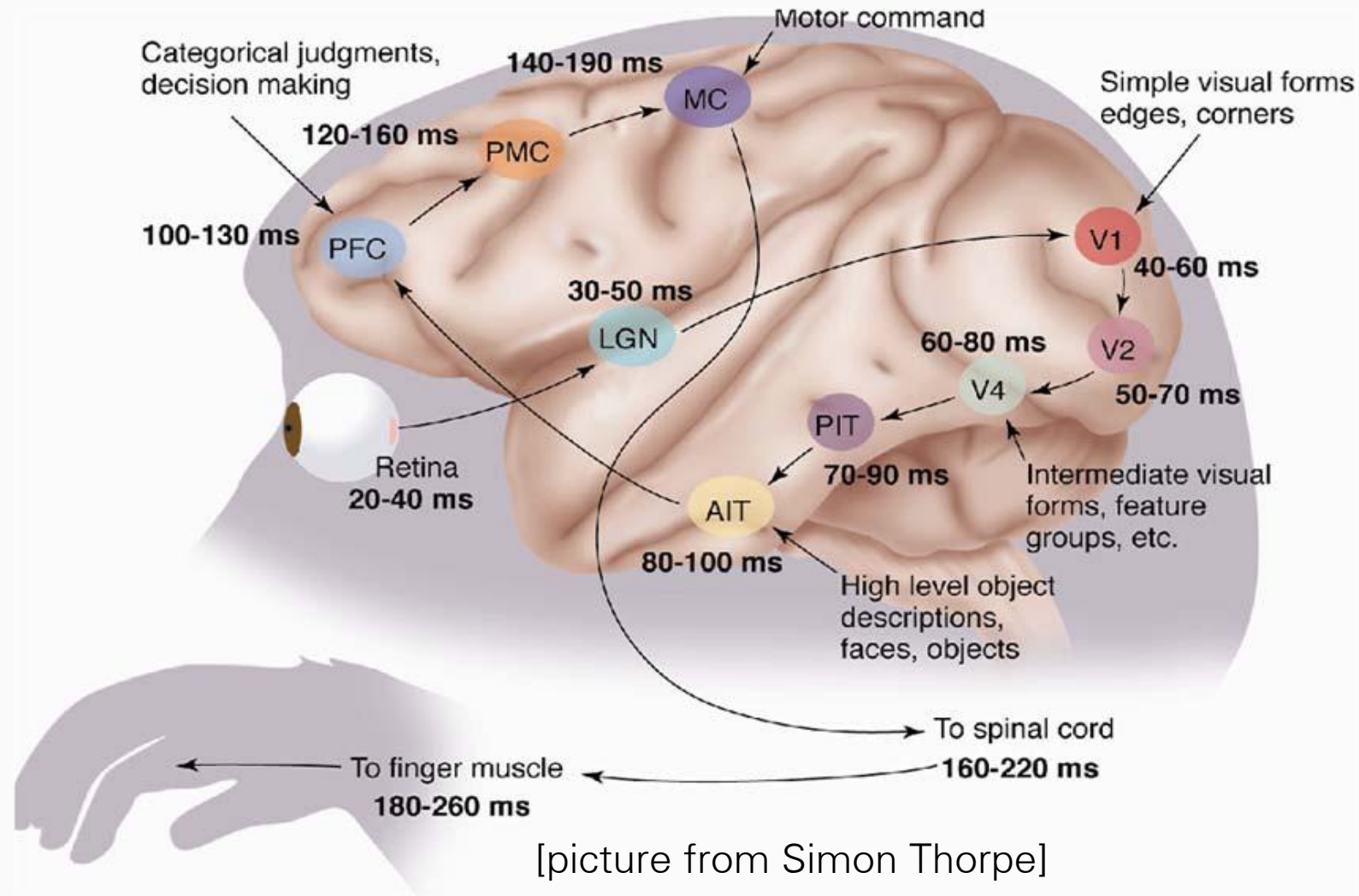


Deep Learning = Hierarchical Compositionality



The Mammalian Visual Cortex is Hierarchical

- The ventral (recognition) pathway in the visual cortex

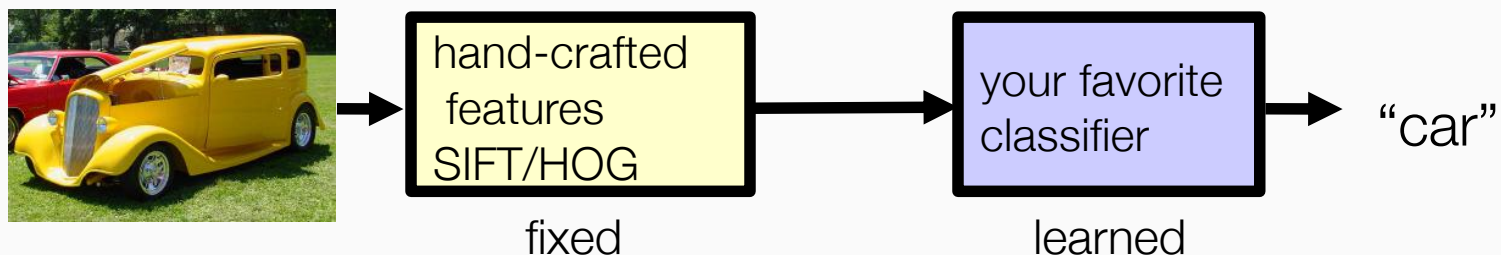


Three key ideas

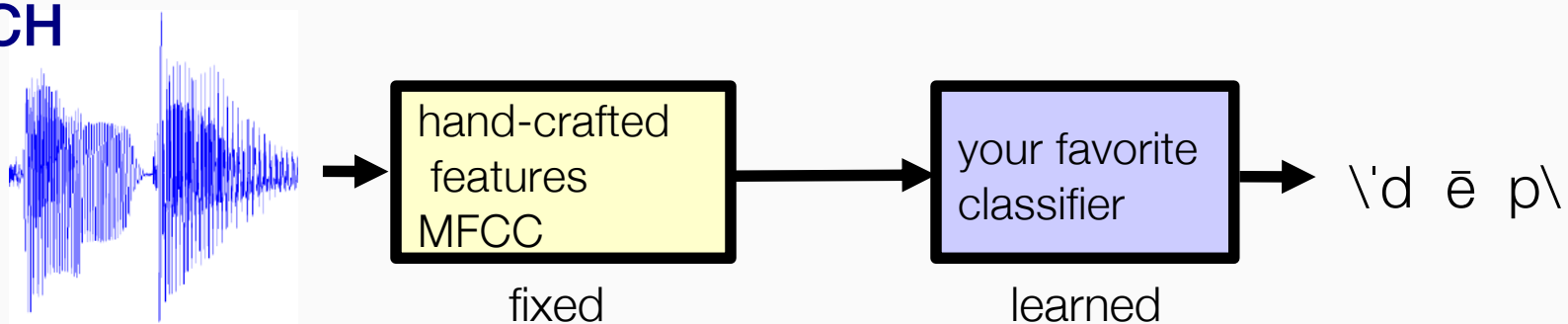
- (Hierarchical) Compositionality
 - Cascade of non-linear transformations
 - Multiple layers of representations
- **End-to-End Learning**
 - Learning (goal-driven) representations
 - Learning to feature extract
- Distributed Representations
 - No single neuron “encodes” everything
 - Groups of neurons work together

Traditional Machine Learning

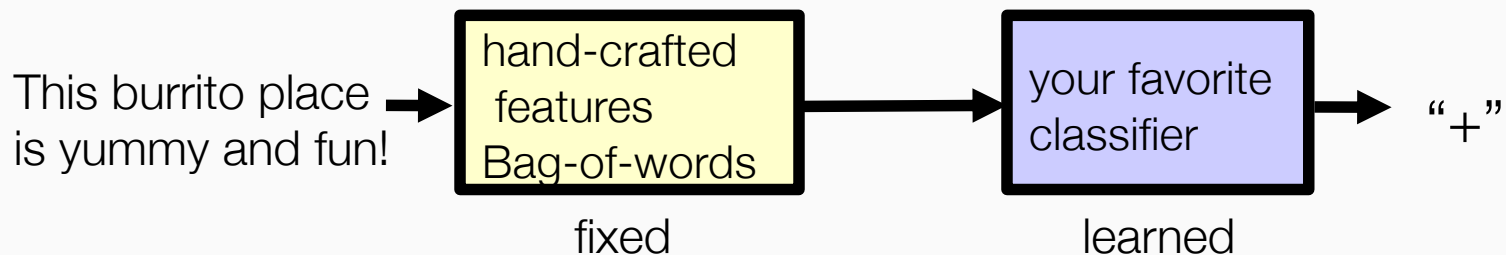
VISION



SPEECH

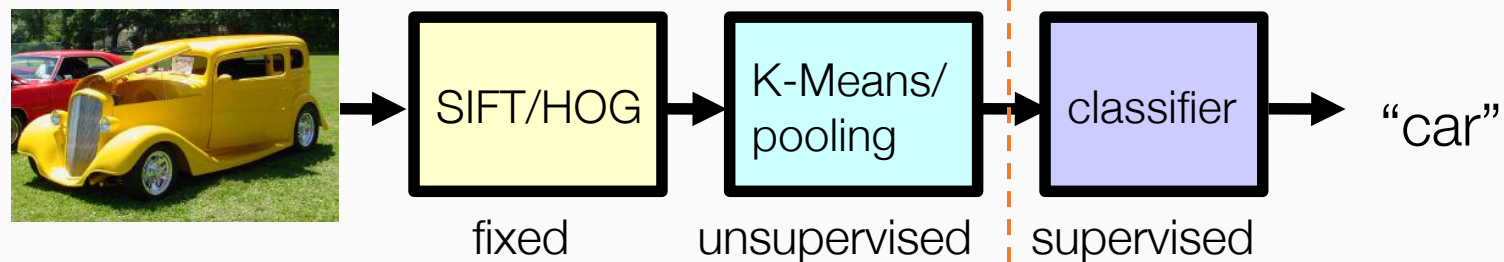


NLP

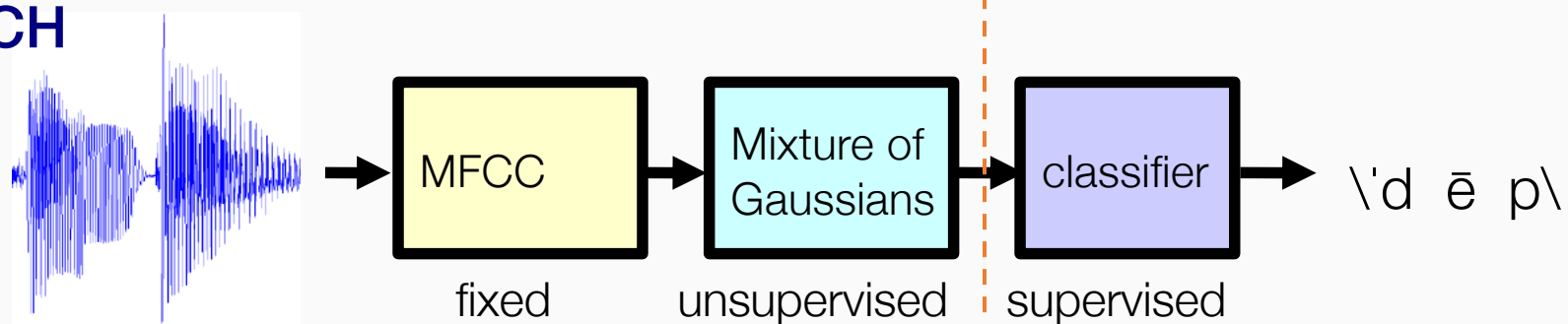


More accurate version

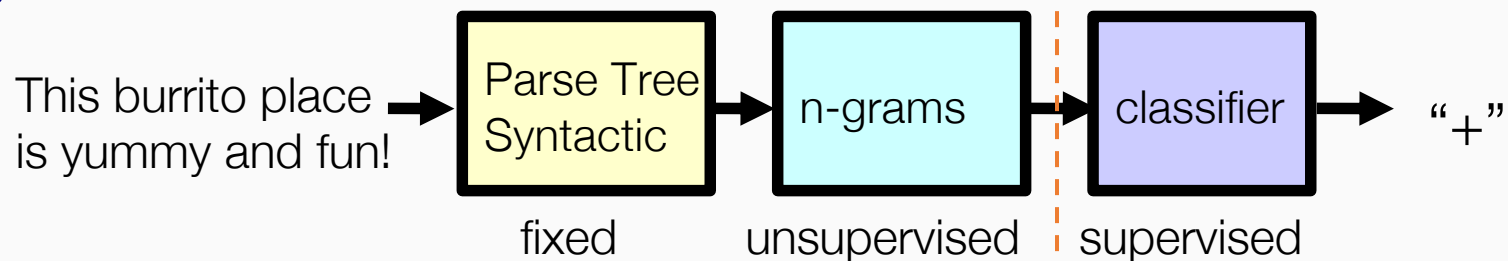
VISION



SPEECH

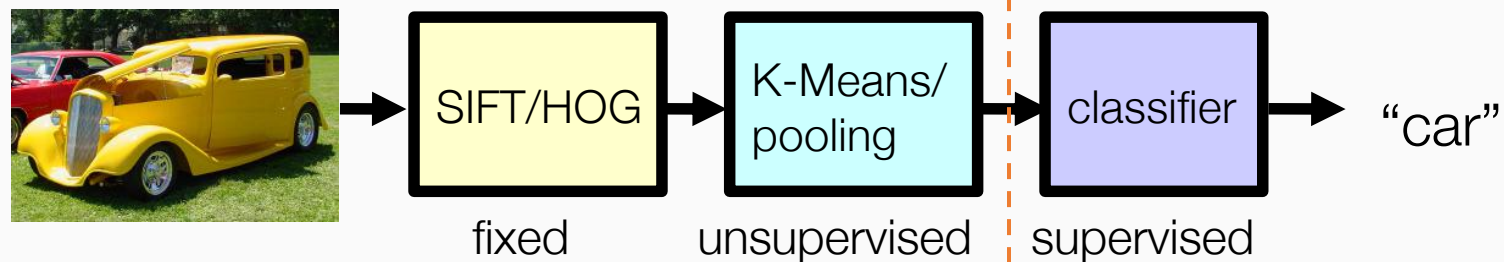


NLP

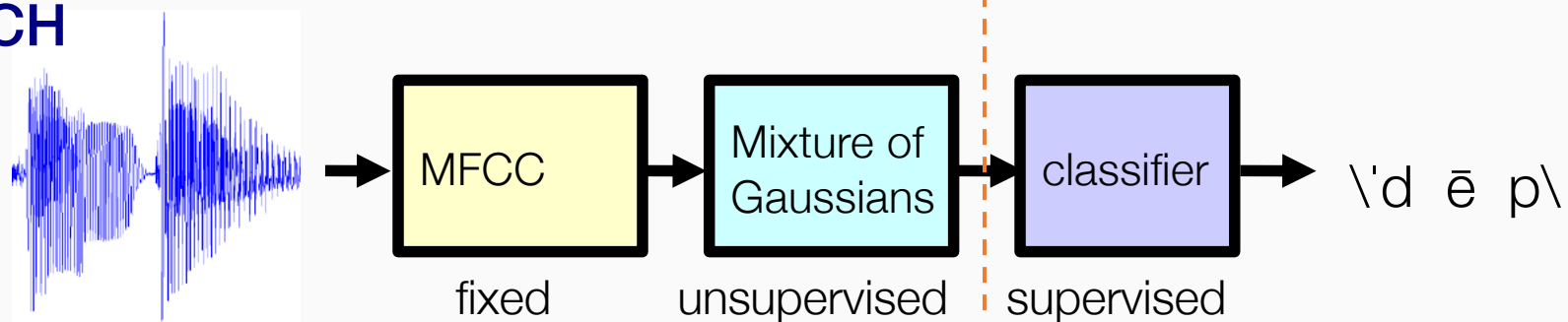


Deep Learning = End-to-End Learning

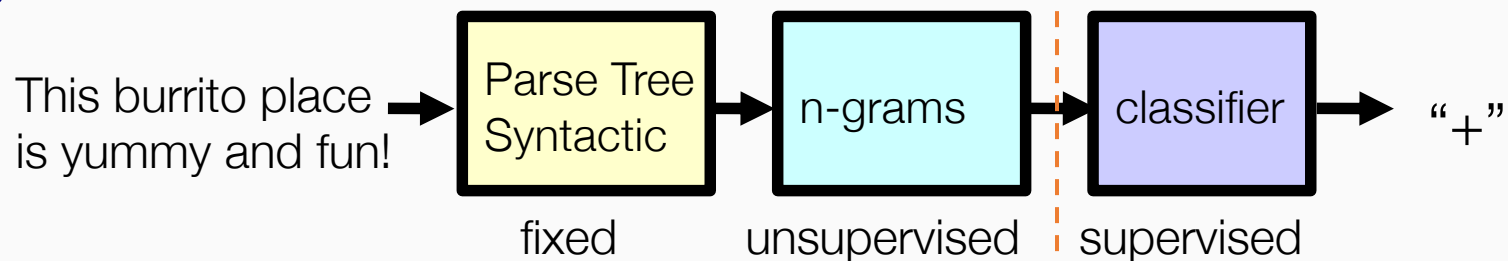
VISION



SPEECH

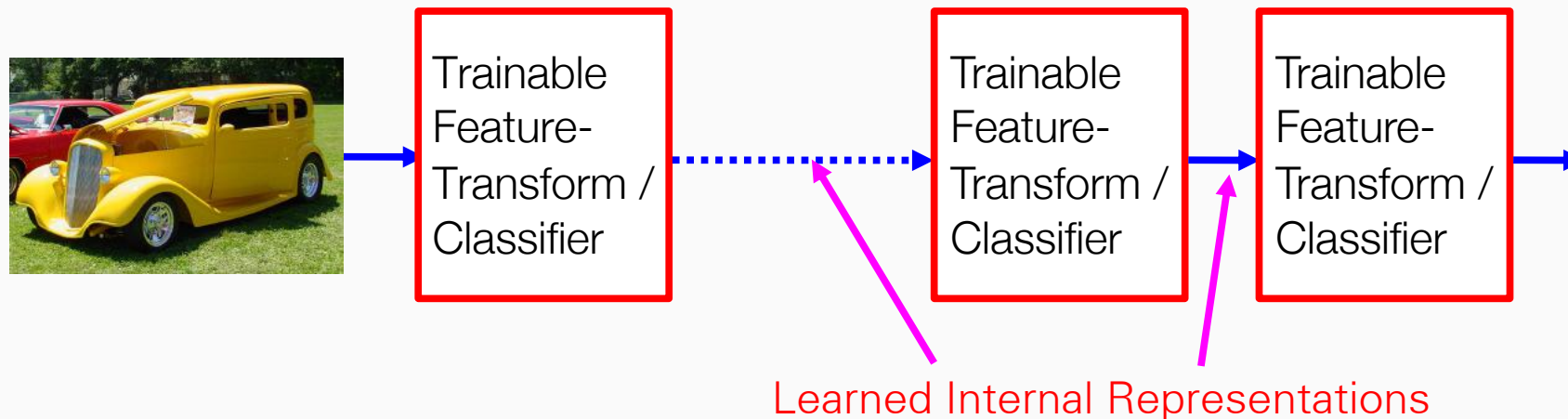


NLP



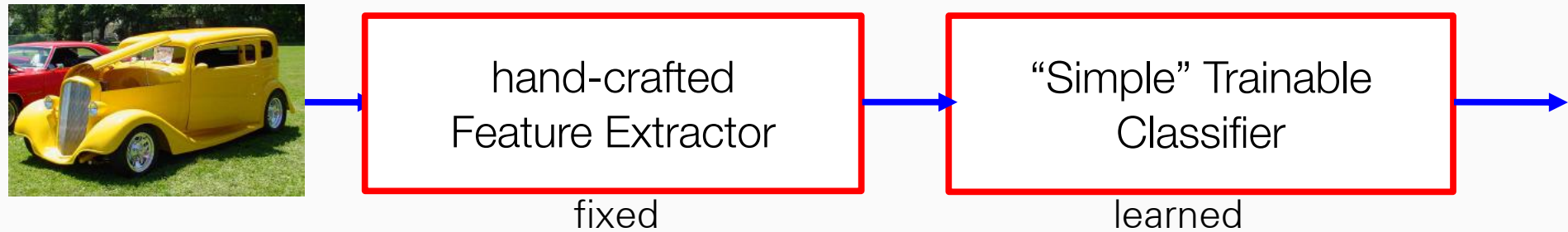
Deep Learning = End-to-End Learning

- A hierarchy of trainable feature transforms
 - Each module transforms its input representation into a higher-level one.
 - High-level features are more global and more invariant
 - Low-level features are shared among categories

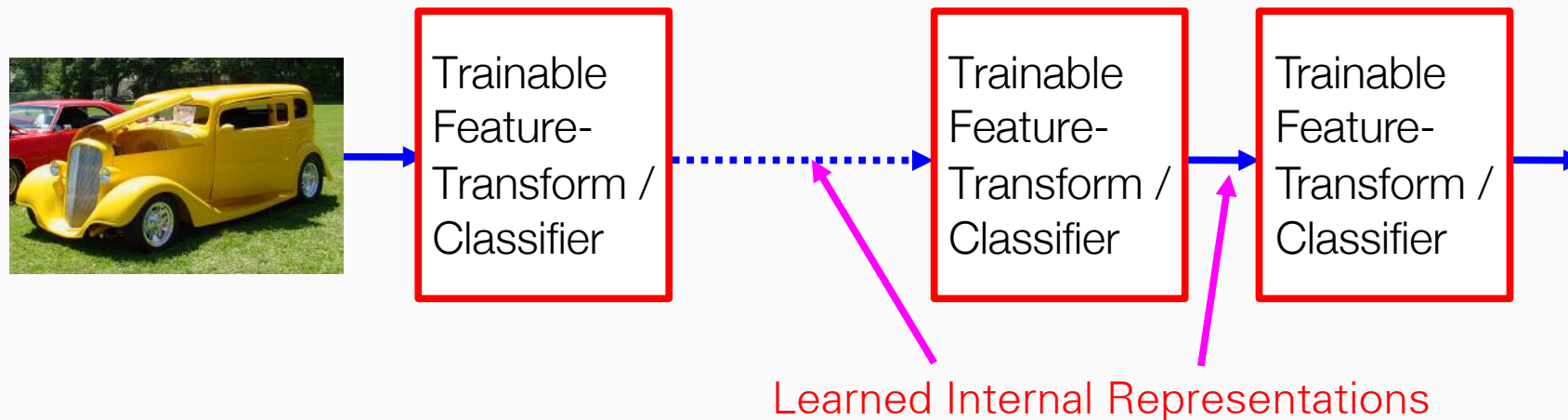


"Shallow" vs Deep Learning

- "Shallow" models



- Deep models

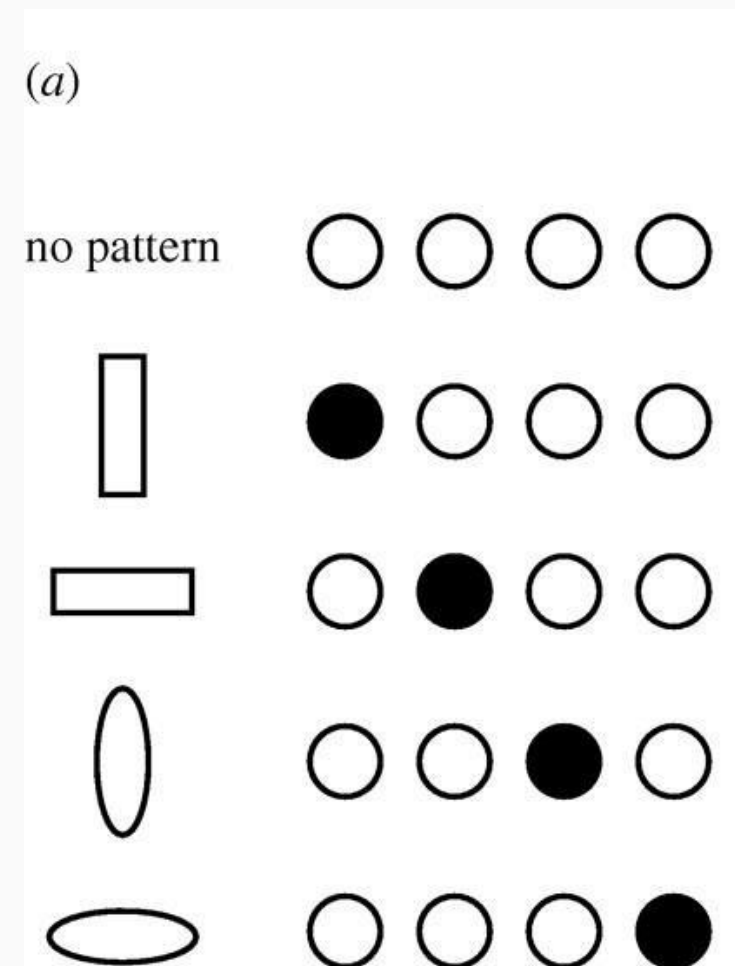


Three key ideas

- (Hierarchical) Compositionality
 - Cascade of non-linear transformations
 - Multiple layers of representations
- End-to-End Learning
 - Learning (goal-driven) representations
 - Learning to feature extract
- **Distributed Representations**
 - No single neuron “encodes” everything
 - Groups of neurons work together

Localist representations

- The simplest way to represent things with neural networks is to **dedicate one neuron to each thing**.
 - Easy to understand.
 - Easy to code by hand
 - Often used to represent inputs to a net
 - Easy to learn
 - This is what mixture models do.
 - Each cluster corresponds to one neuron
 - Easy to associate with other representations or responses.
- But localist models are very inefficient whenever the data has componential structure.

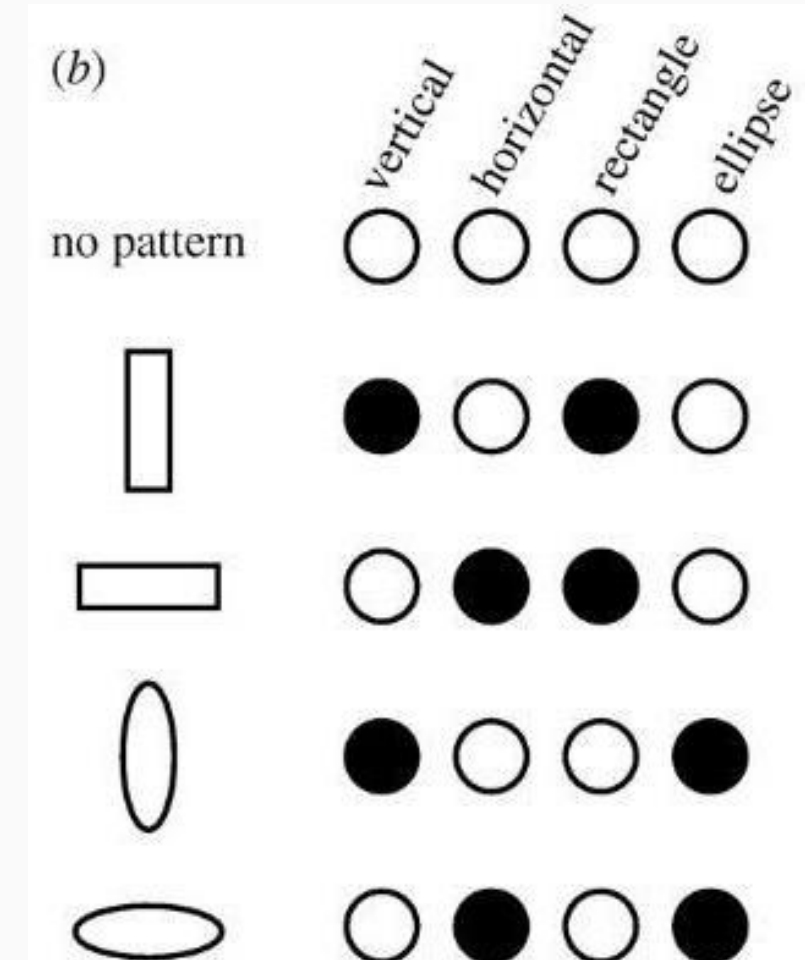


Distributed Representations

- Each neuron must represent something, so this must be a local representation.
- **Distributed representation** means a many-to-many relationship between two types of representation (such as concepts and neurons).
 - Each concept is represented by many neurons
 - Each neuron participates in the representation of many concepts

Local ● ● ○ ● = VR + HR + HE = ?

Distributed ● ● ○ ● = V + H + E ≈ ○



Power of distributed representations!

Scene Classification



- Possible internal representations:

- Objects
- Scene attributes
- Object parts
- Textures



Simple elements & colors

Object part

Object

Scene

Three key ideas of deep learning

- **(Hierarchical) Compositionality**

- Cascade of non-linear transformations
- Multiple layers of representations

- **End-to-End Learning**

- Learning (goal-driven) representations
- Learning to feature extract

- **Distributed Representations**

- No single neuron “encodes” everything
- Groups of neurons work together

Benefits of Deep/Representation Learning

- (Usually) Better Performance
 - “Because gradient descent is better than you”
Yann LeCun
- New domains without “experts”
 - RGBD
 - Multi-spectral data
 - Gene-expression data
 - Unclear how to hand-engineer

Problems with Deep Learning

- **Problem#1: Non-Convex! Non-Convex! Non-Convex!**

- Depth ≥ 3 : most losses non-convex in parameters
- Theoretically, all bets are off
- Leads to stochasticity
 - different initializations $\rightarrow$ different local minima

- Standard response #1

- “Yes, but all interesting learning problems are non-convex”
- For example, human learning
 - Order matters $\rightarrow$ wave hands $\rightarrow$ non-convexity

- Standard response #2

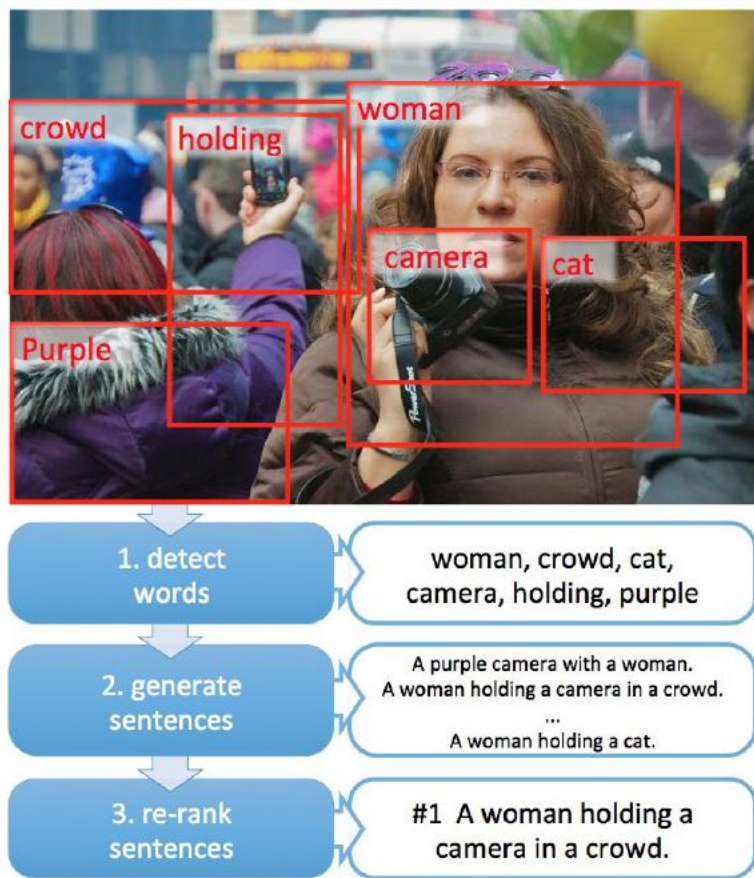
- “Yes, but it often works!”

Problems with Deep Learning

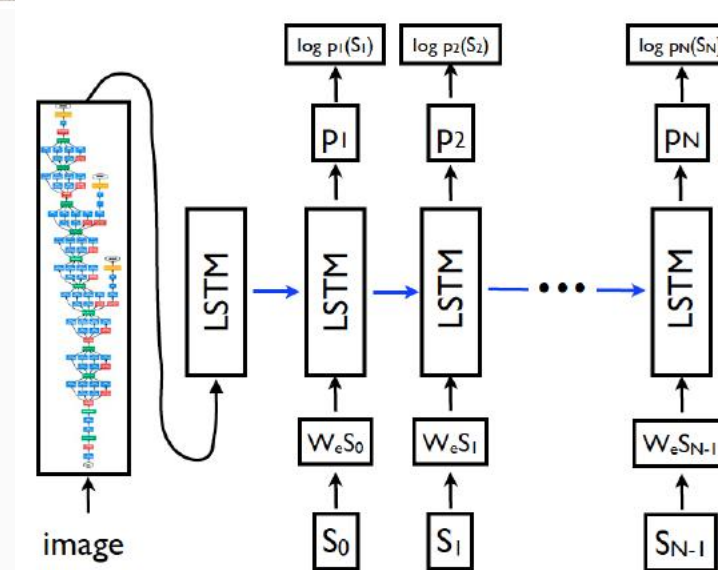
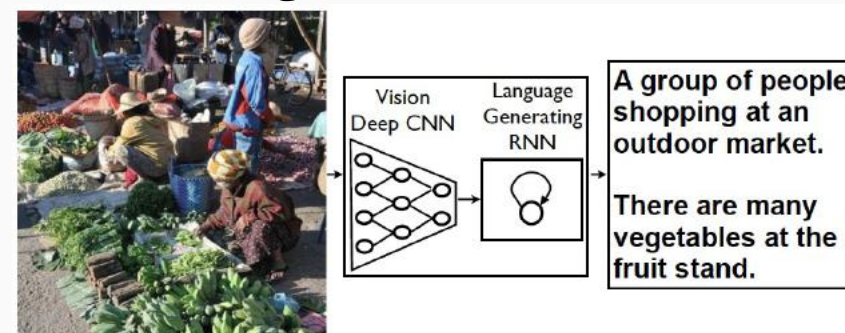
- **Problem#2: Hard to track down what's failing**
 - Pipeline systems have “oracle” performances at each step
 - In end-to-end systems, it's hard to know why things are not working

Problems with Deep Learning

- Problem#2: Hard to track down what's failing



[Fang et al. CVPR15]



[Vinyals et al. CVPR15]

Pipeline

End-to-End

Problems with Deep Learning

- **Problem#2: Hard to track down what's failing**
 - Pipeline systems have “oracle” performances at each step
 - In end-to-end systems, it's hard to know why things are not working
- Standard response #1
 - Tricks of the trade: visualize features, add losses at different layers, pre-train to avoid degenerate initializations...
 - “We're working on it”
- Standard response #2
 - “Yes, but it often works!”

Problems with Deep Learning

- **Problem#3: Lack of easy reproducibility**
 - Direct consequence of stochasticity & non-convexity
- Standard response #1
 - It's getting much better
 - Standard toolkits/libraries/frameworks now available
- Standard response #2
 - “Yes, but it often works!”

NEW NAVY DEVICE LEARNS BY DOING

Psychologist Shows Embryo
of Computer Designed to
Read and Grow Wiser

WASHINGTON, July 7 (UPI)—The Navy revealed the embryo of an electronic computer today that it expects will be able to walk, talk, see, write, reproduce itself and be conscious of its existence.

The embryo—the Weather Bureau's \$2,000,000 "704" computer—learned to differentiate between right and left after fifty attempts in the Navy's demonstration for newsmen.

The service said it would use this principle to build the first of its Perceptron thinking machines that will be able to read and write. It is expected to be finished in about a year at a cost of \$100,000.

Dr. Frank Rosenblatt, designer of the Perceptron, conducted the demonstration. He said the machine would be the first device to think as the human brain. As do human be-

ings, Perceptron will make mistakes at first, but will grow wiser as it gains experience, he said.

Dr. Rosenblatt, a research psychologist at the Cornell Aeronautical Laboratory, Buffalo, said Perceptrons might be fired to the planets as mechanical space explorers.

Without Human Controls

The Navy said the perceptron would be the first non-living mechanism "capable of receiving, recognizing and identifying its surroundings without any human training or control."

The "brain" is designed to remember images and information it has perceived itself. Ordinary computers remember only what is fed into them on punch cards or magnetic tape.

Later Perceptrons will be able to recognize people and call out their names and instantly translate speech in one language to speech or writing in another language, it was predicted.

Mr. Rosenblatt said in principle it would be possible to build brains that could reproduce themselves on an assembly line and which would be conscious of their existence.

1958 New York
Times...

In today's demonstration, the "704" was fed two cards, one with squares marked on the left side and the other with squares on the right side.

Learns by Doing

In the first fifty trials, the machine made no distinction between them. It then started registering a "Q" for the left squares and "O" for the right squares.

Dr. Rosenblatt said he could explain why the machine learned only in highly technical terms. But he said the computer had undergone a "self-induced change in the wiring diagram."

The first Perceptron will have about 1,000 electronic "association cells" receiving electrical impulses from an eye-like scanning device with 400 photo-cells. The human brain has 10,000,000,000 responsive cells, including 100,000,000 connections with the eyes.

The New York Times

Science

WORLD

U.S.

N.Y. / REGION

BUSINESS

TECHNOLOGY

SCIENCE

HEALTH

SPORTS

OPINION

ENVIRONMENT SPACE & COSMOS

COMPUTER SCIENTISTS STYMIED IN THEIR QUEST TO MATCH HUMAN VISION

By WILLIAM J. BROAD

Published: September 25, 1984

EXPERTS pursuing one of man's most audacious dreams - to create machines that think - have stumbled while taking what seemed to be an elementary first step. They have failed to master vision.

After two decades of research, they have yet to teach machines the seemingly simple act of being able to recognize everyday objects and to distinguish one from another.

Instead, they have developed a profound new respect for the sophistication of human sight and have scoured such fields as mathematics, physics, biology and psychology for clues to help them achieve the goal of machine vision.



FACEBOOK



TWITTER



GOOGLE+



EMAIL



SHARE



PRINT



REPRINTS

SCIENCE

Researchers Announce Advance in Image-Recognition Software

By JOHN MARKOFF NOV. 17, 2014



Email



Share



Tweet



Save



More

MOUNTAIN VIEW, Calif. — Two groups of scientists, working independently, have created artificial intelligence software capable of recognizing and describing the content of photographs and videos with far greater accuracy than ever before, sometimes even mimicking human levels of understanding.

Until now, so-called computer vision has largely been limited to recognizing individual objects. The new software, described on Monday by researchers at Google and at [Stanford University](#), teaches itself to identify entire scenes: a group of young men playing Frisbee, for example, or a herd of elephants marching on a grassy plain.

The software then writes a caption in English describing the picture. Compared with human observations, the researchers found, the computer-written descriptions are surprisingly accurate.

Captioned by Human and by Google's Experimental Program



Human: "A group of men playing Frisbee in the park."

Computer model: "A group of young people playing a game of Frisbee."

TWEETS
587

FOLLOWING
18

FOLLOWERS
746

FAVORITES
13



INTERESTING.JPG @INTERESTING_JPG · 10h

a man holding a mirror up to his face .



[View more photos and videos](#)

Results from @INTERESTING_JPG via <http://deeplearning.cs.toronto.edu/i2t>

TWEETS
587

FOLLOWING
18

FOLLOWERS
746

FAVORITES
13



INTERESTING.JPG @INTERESTING_JPG · 18h

a man carrying a bucket of his hands in a yard .



[View more photos and videos](#)

Results from @INTERESTING_JPG via <http://deeplearning.cs.toronto.edu/i2t>

TWEETS
587

FOLLOWING
18

FOLLOWERS
746

FAVORITES
13



INTERESTING.JPG @INTERESTING_JPG · Feb 20

a surfboard attached to the top of a car .



[View more photos and videos](#)

TWEETS
587

FOLLOWING
18

FOLLOWERS
746

FAVORITES
13



INTERESTING.JPG @INTERESTING_JPG · Feb 19

a man dressed in uniform is looking at his cell phone .



[View more photos and videos](#)

Results from @INTERESTING_JPG via <http://deeplearning.cs.toronto.edu/i2t>

TWEETS
587

FOLLOWING
18

FOLLOWERS
746

FAVORITES
13



INTERESTING.JPG @INTERESTING_JPG · 16h

this appears to be a small bedroom in the snow .



[View more photos and videos](#)

Results from @INTERESTING_JPG via <http://deeplearning.cs.toronto.edu/i2t>



Iain Murray

@driainmurray

Follow

Today I learned [#googletranslate](#) sometimes decides that "Deutsch" means "English". Machine learning systems need to cope with weird inputs.

Translate

Turn off instant translation

RussianGermanEnglishDetect language

EnglishGermanSpanish

Translate

Deutschland

Deutsch, deutsch, deutsch, deutsch, deutsch, deutsch

Natürlich hat ein Deutscher "Wetten, dass ...?" erfunden
Vielen Dank für die schönen Stunden!
Wir sind die freundlichsten Kunden auf dieser Welt
Wir sind bescheiden, wir haben Geld
Die Allerbesten in jedem Sport
Die Steuern hier sind Weltrekord
Bereisen Sie Deutschland und bleiben Sie hier!
Auf diese Art von Besuchern warten wir
Es kann jeder hier wohnen, dem es gefällt
Wir sind das freundlichste Volk auf dieser Welt

Deutsch, deutsch, deutsch, deutsch

Germany

German, English, German, German, German, and English

Of course a German has "betting that ...?" invented
Thanks for the nice hours!
We are the friendliest customers in this world
We are modest, we have money
The very best in any sport
The taxes here are a world record
Travel to Germany and stay here!
We are waiting for this kind of visitors
Anyone who likes it can live here
We are the friendliest people in this world

English, German, German, and German

#OzgurOrman
2,516 Tweets

#TurkeySaysYes
1,520 Tweets

#BahisSarayındaKazandım

Igor Tudor

5,727 Tweets

#valentines

@TEDTalks, @MIT and 5 more are Tweeting about this

#Yellen

2,287 Tweets

© 2017 Twitter About Help Center Terms
Privacy Cookies Ads info



Iain Murray

@driainmurray

Academic in Machine Learning and Statistics.

homepages.inf.ed.ac.uk/imurray2/

Joined May 2011



Iain Murray

@driainmurray

 Follow



More fun pushing [#googletranslate](#)'s neural net into weird states. (BTW try GT on real text if you haven't recently. It's often amazing.)

English German Spanish Detect language

↔

German English Spanish

Translate

knife, fork, knife,

(The trailing comma messes this one up.)

19/5000

Messer, Messer, Messer,

English German Spanish Detect language

↔

German English Spanish

Translate

Messer, Gabel, Messer, Messer, Messer,

Messer, Messer, Messer, Messer, Messer

77/5000

Screen monitor styling Projector styling Print styling ← back to. 2010-01-20 with adjustable interlinear. Knife, fork; knife, knife, knife, knife;

RETWEETS

120

LIKES

184





Tomer Ullman
@TomerUllman



Do models like DALL-E 2 get basic relations (in/on/etc)?

Colin (Coco) Conwell and I set out to investigate. The result is now on arXiv:

“Testing Relational Understanding in Text-Guided Image Generation”



arxiv.org

Testing Relational Understanding in Text-Guided Image Gen...
Relations are basic building blocks of human cognition.
Classic and recent work suggests that many relations are ...

2:55 PM · Aug 2, 2022 · Twitter Web App

“A spoon in a cup”



“A cup on a spoon”





Melanie Mitchell

@MelMitchell1



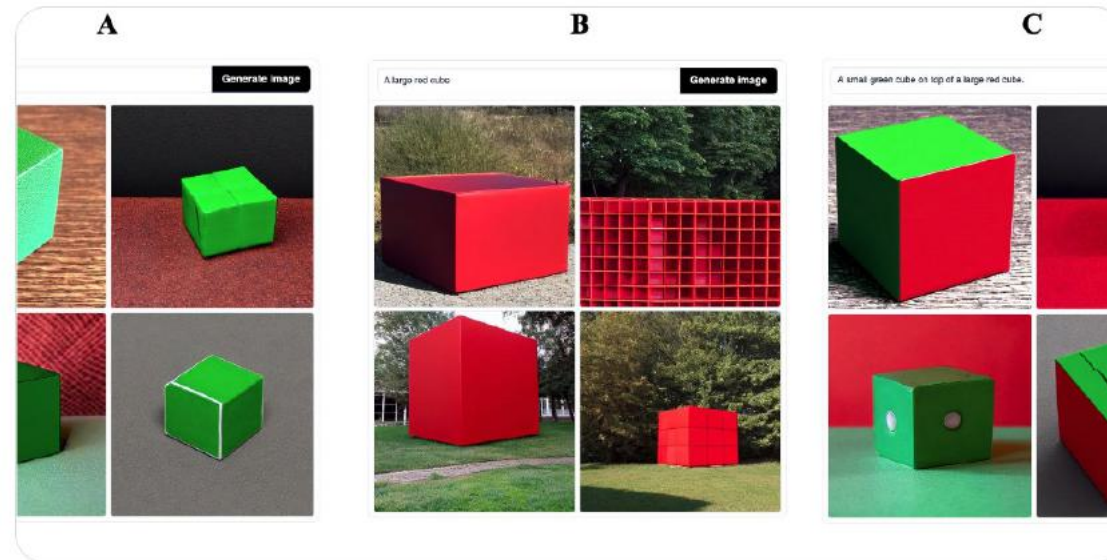
Prepositions are hard.

Stable diffusion demo ([huggingface.co/spaces/stabili...](https://huggingface.co/spaces/stabilityai/stable-diffusion))

Prompt A: A small green cube

Prompt B: A large red cube

Prompt C: A small green cube on top of a large red cube



6:10 PM · Aug 23, 2022 · Twitter Web App



Melanie Mitchell
@MelMitchell1

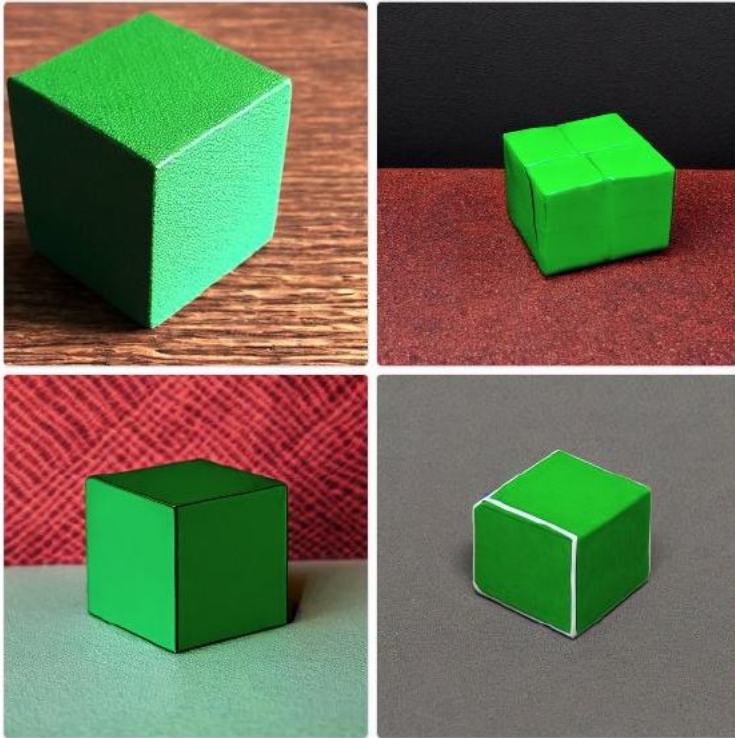


Propositions are hard

A

A small green cube

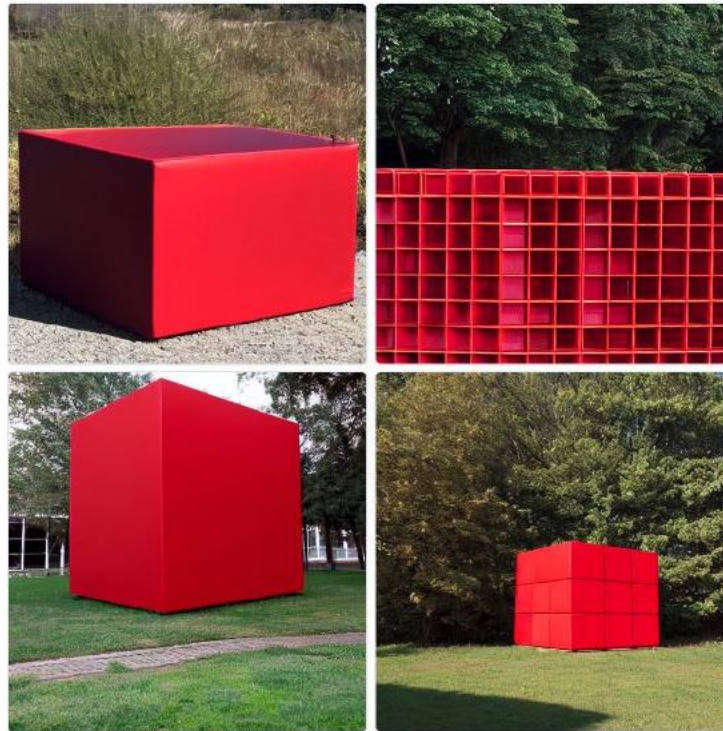
Generate image



B

A large red cube

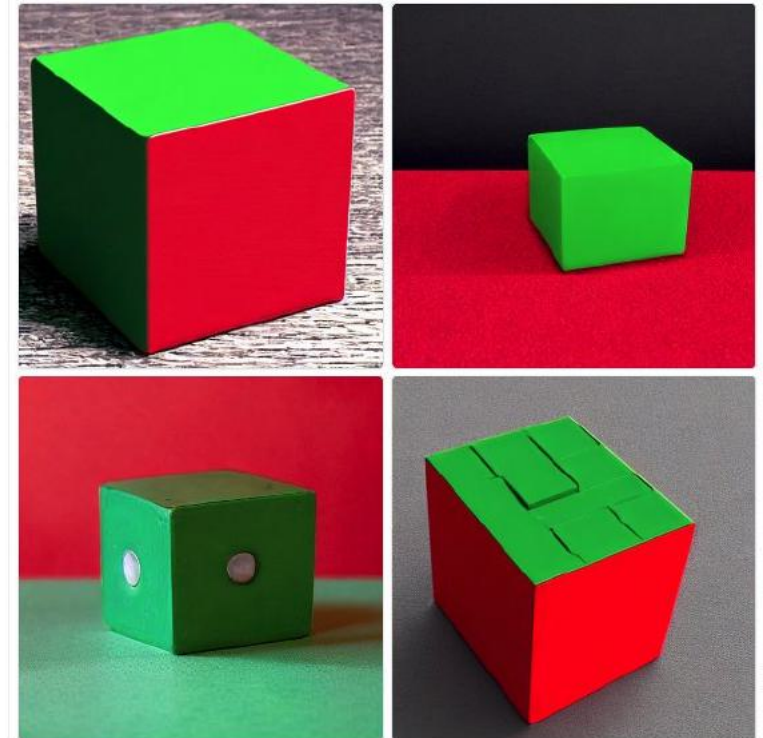
Generate image



C

A small green cube on top of a large red cube.

Generate image



6:10 PM · Aug 23, 2022 · Twitter Web App



Melanie Mitchell
@MelMitchell1



**Propositions are hard **

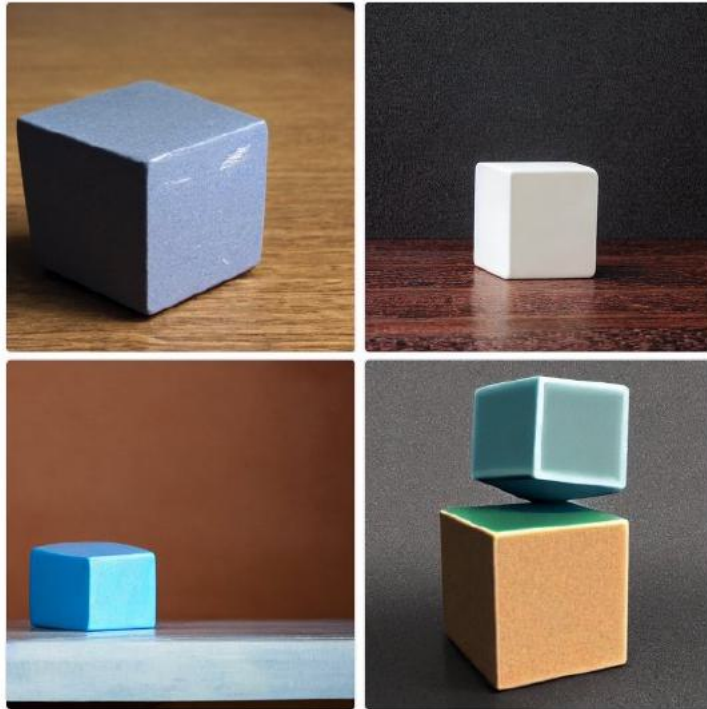
A

B

C

One cube on top of another cube

Generate image



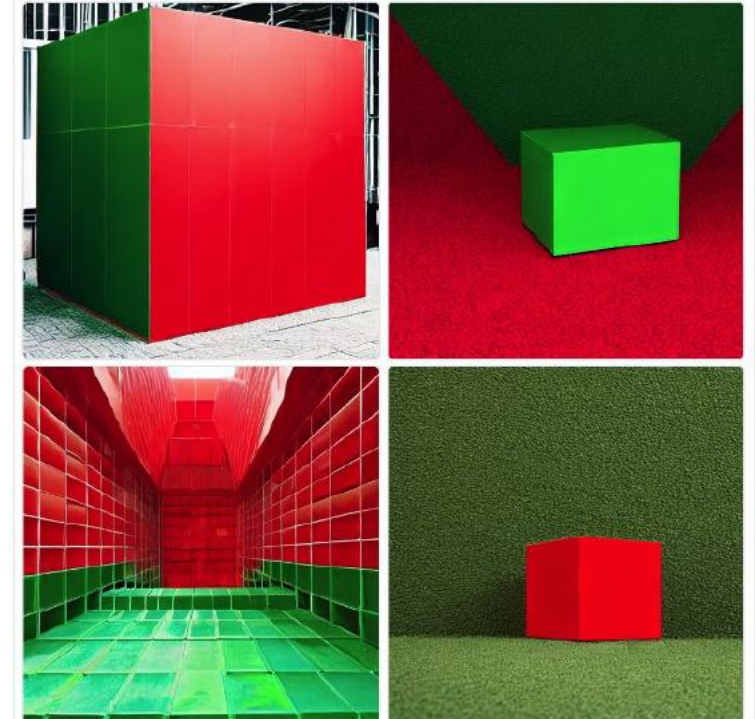
A small cube to the left of a large cube

Generate image



A red cube below a green cube

Generate Image



6:10 PM · Aug 23, 2022 · Twitter Web App



Gary Marcus ✓
@GaryMarcus



The disembodied hand is a nice touch.

“a donkey rides on top of a bass guitarist who is on rollerskates, as they cross a football field. a mouse sits on top of the donkey.”



4:10 PM · Oct 5, 2025 · **67K** Views



Gary Marcus ✓
@GaryMarcus

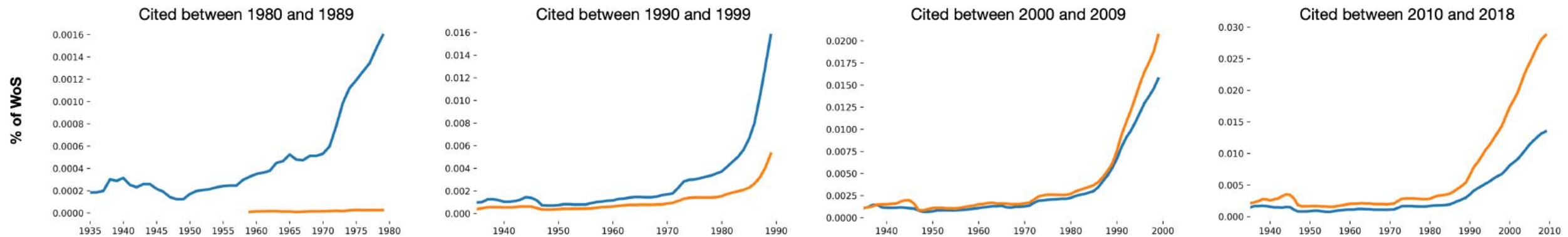
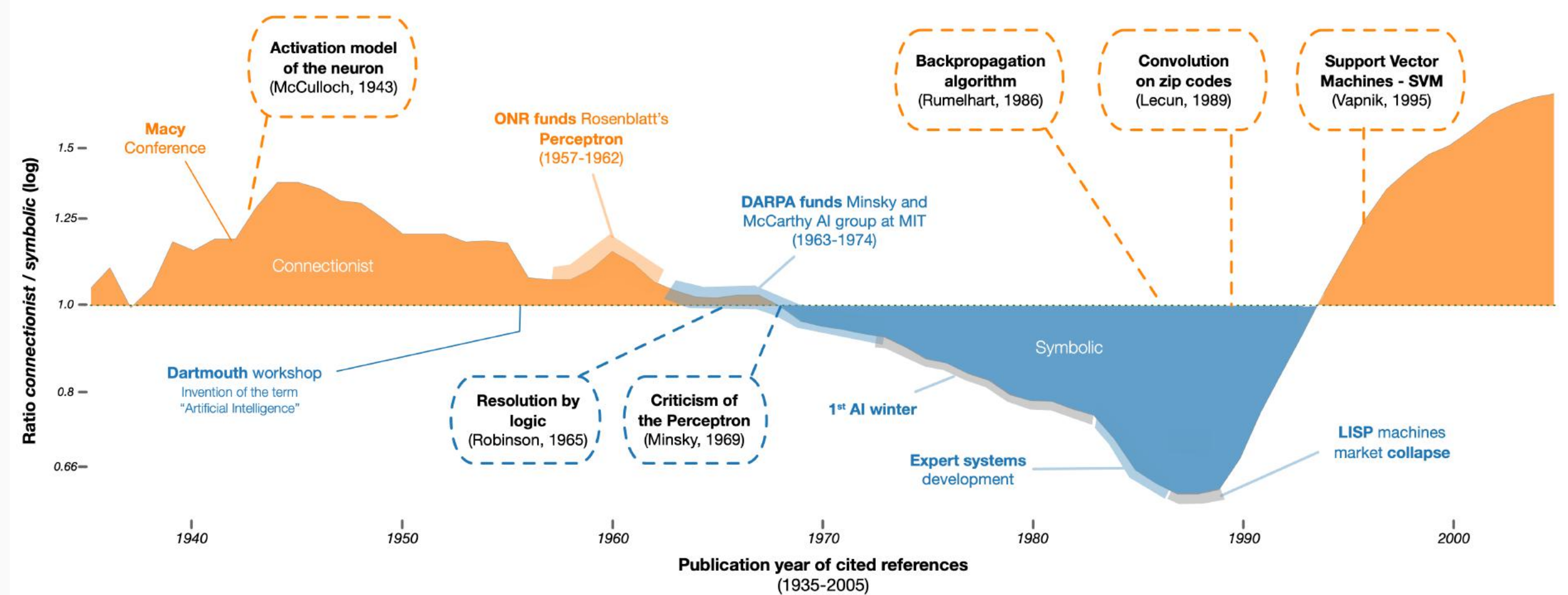


The disembodied hand is a nice touch.

“a donkey rides on top of a bass guitarist who is on rollerskates, as they cross a football field. a mouse sits on top of the donkey.”



4:10 PM · Oct 5, 2025 · **67K** Views



Next Lecture: Machine Learning Overview