

EVI A2025



USE CASE: IMAGE RECOGNITION

Petia Radeva, IAPR Fellow, ELLIS Member

Full professor

Head of “Artificial Intelligence and Biomedical
Applications” Consolidated Research Group,

Universitat de Barcelona, Spain &

Barcelona Supercomputing Center

University of Barcelona

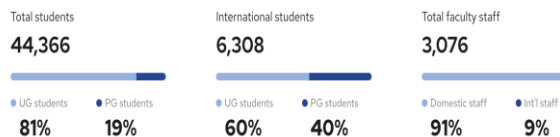


In 1450, King Alfonso V the Magnanimous established the University of Barcelona.

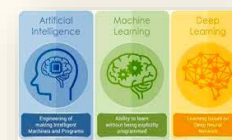
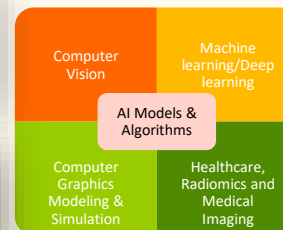
University information

STUDENT & STAFF

16 Faculties



- Ranked 1st Spanish university in 2023-2024 rankings
 - THE World University Rankings, (#152)
 - SCIMAGO Institutions ranking 2024 (#91)
 - QS World University Rankings 2025 (#165)



Consolidated research group "Artificial Intelligence and Biomedical Applications"

WHO AM I?



Petia Radeva



Full professor at the
University of Barcelona,
Dept. Mathematics and
Informatics,

Consolidated Research Group
**Artificial Intelligence and Bio-
Medical Applications**

Head of AIBA



Editor in Chief of
Pattern Recognition,
Q1, IF: 7.6.



Google scholar h-
index is 59 with
>12000 cites



280 journals
and chapters;
20 scientific
books



80+ invited talks in conferences
(MetaFood CVPRW'24) and
summer schools



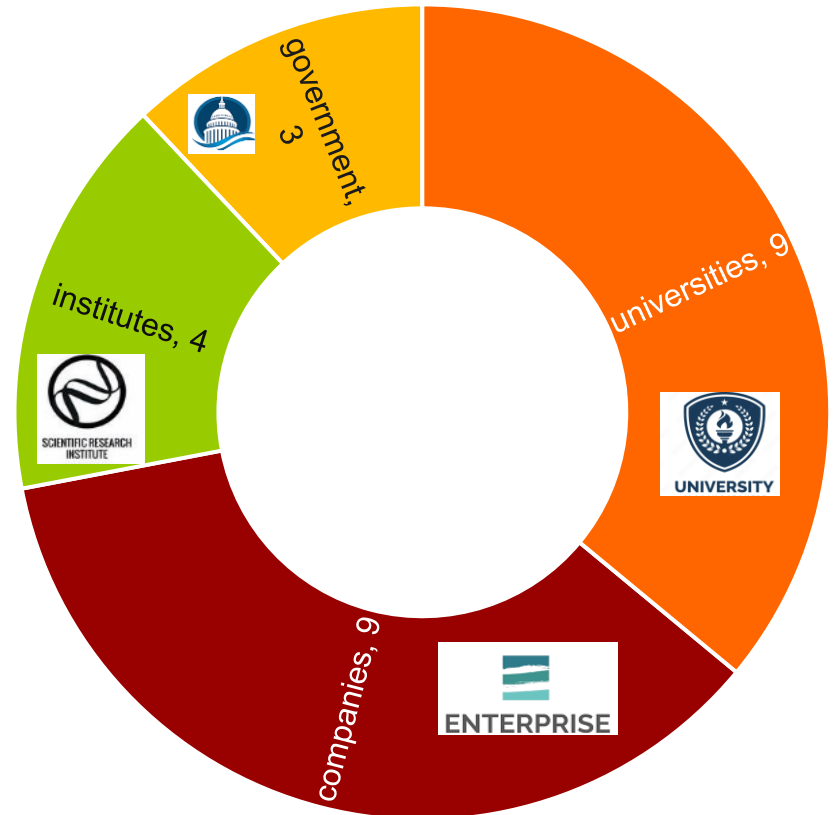
- Program Co-Chair of ACM Multimedia 2027, Hong Kong; Area chair of CVPR'2025, WACV'2024, ICCV'2023, CVPR'2022, Track chair of ICPR'2024, Program chair of IBPRIA'25, VISAPP'24, VISAPP'23, VISAPP'22, VISAP'21, VISAPP'20, etc.

1

Awards: “Narcis Monturiol” medal, ICT Research Woman'2025, IAPR Fellow, “Aurora Pons” CIARP Award, 1 of 100 ForbesWoman'2025, etc.



What am I most proud of?

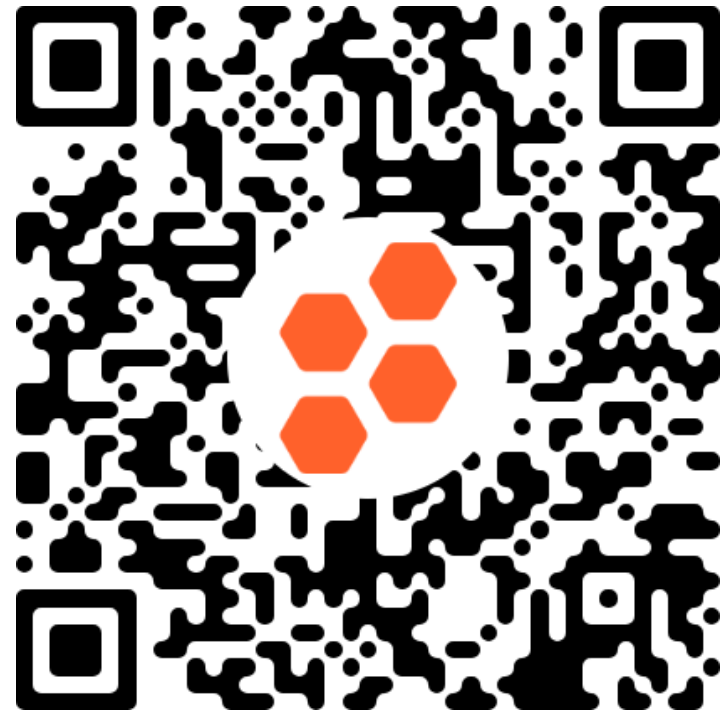


Let's know us

Let's know me
your opinion

Go to
goSocrative.com

Enter to room:
petiaradeva



What are we doing?



Google's DeepMind A.I. beats doctors in breast cancer screening trial

PUBLISHED THU, JAN 2 2020 8:10 AM EST | UPDATED THU, JAN 2 2020 8:10 AM EST

David Nield

KEY POINTS

- Anonymous scans of 29,000 women were used in the trial.
- The biggest improvements over human scanning was found in the U.S. portion of the study.
- Google-owned DeepMind has already used AI to read eye scans and spot neck cancer.



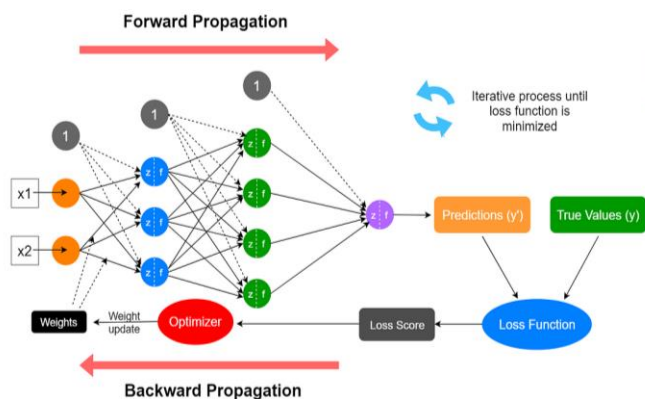
Become part of the Revolution

Get trained, get connected and join by the most successful graduate employer

AmplifyME

Learn more >

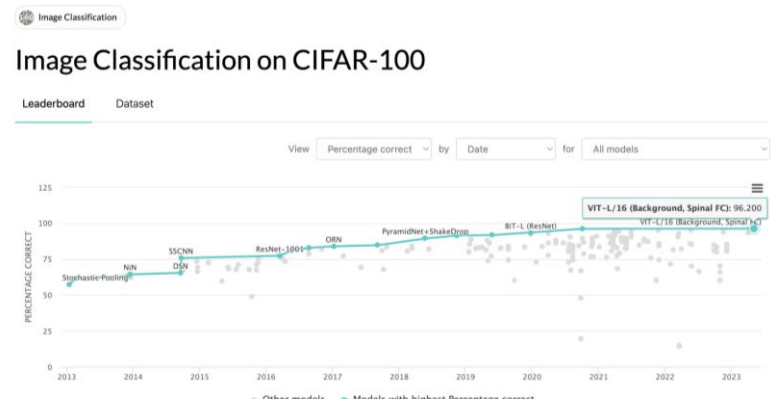
TRENDING NOW



© 2011 Pearson Education, Inc. or its affiliate(s). All rights reserved. 10

HOW MANY OF YOU

- ❖ Are using annotated datasets?
- ❖ Are (not) using annotated datasets?
- ❖ Did you know that:
 - ❖ 3.4% average error rate across all datasets, 6% for ImageNet
 - ❖ MNIST digits dataset contains 15 (human-validated) label errors in the test set.



So there are 2 main questions:

- How to find the noisy labels?
- Can we use them to make algorithms more robust?



Sample Selection Techniques

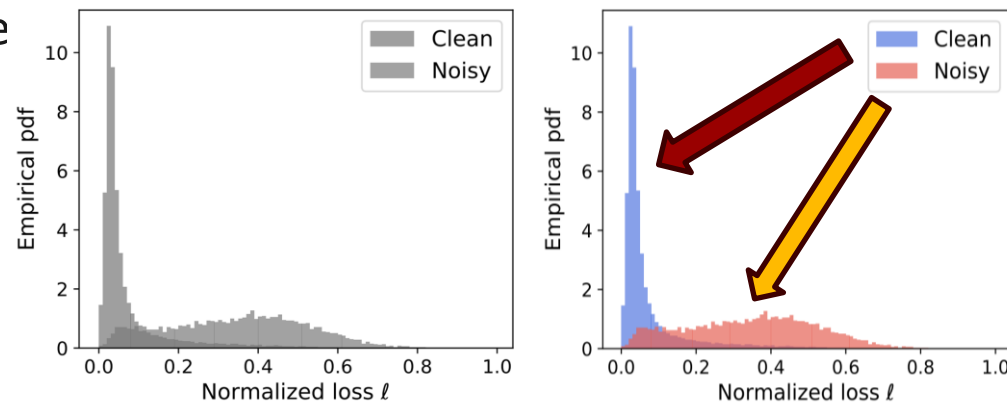
Question: What is a good selection criterion to split clean and noisy data?



Common strategy:

Using Loss Distribution to split the

Small loss $\rightarrow$ Clean Samples



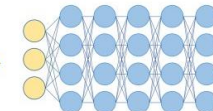
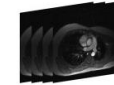
¿What makes NNs so popular?



Advantages:

1. End-to-end learning

DEEP LEARNING

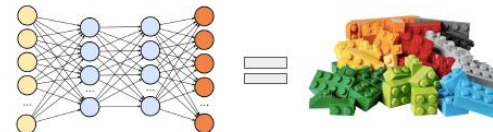


Feature extraction & classification



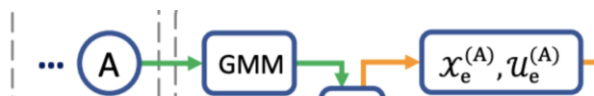
Class 1 / Class 2

2. Modularity



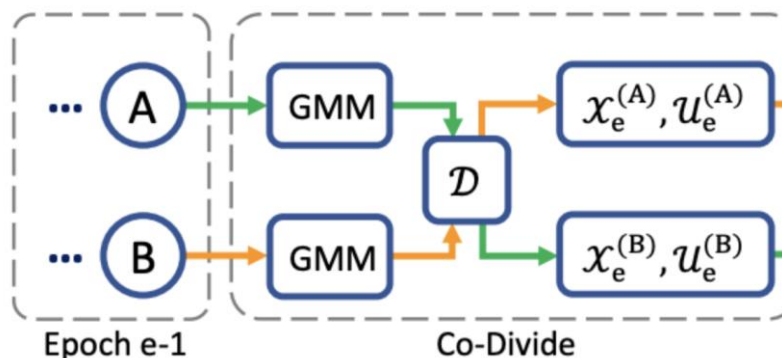
Confirmation Bias

- ❖ Training a model using its own data division causes **confirmation bias**

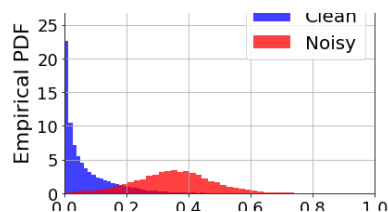
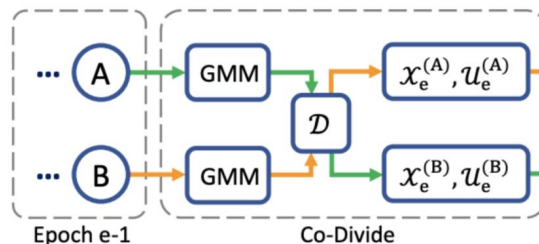


Overfitting

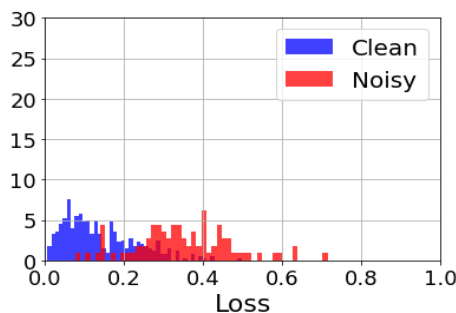
- ❖ Because noisy samples that are **wrongly grouped** into the labelled set would keep **lower loss due to the model overfitting** to the labels.



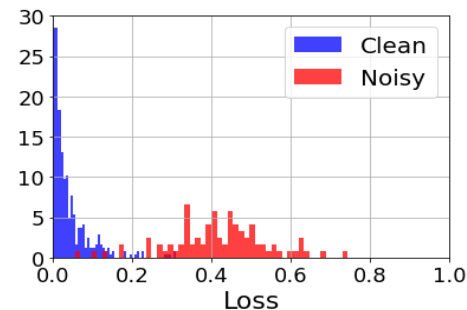
Class-conditional Learning with Noisy Labeling



All classes



Class 'Otter'



Class 'Road'

Loss function

$$\mathcal{L} = \mathcal{L}_{\mathcal{X}} + \lambda_u \mathcal{L}_{\mathcal{U}} + \lambda_r \mathcal{L}_{\text{reg}}$$

Uncertainty

$$\mathcal{L}_{\mathcal{X}} = -\frac{1}{|\mathcal{X}|} \sum_{x,y,p \in \mathcal{X}} (1 + \lambda_{\mathcal{X}} \cdot \psi^y) \cdot p_c \log(p_{\text{model}}^c(x; \theta)),$$

LABELED LOSS

$$\mathcal{L}_{\mathcal{U}} = \frac{1}{|\mathcal{U}|} \sum_{x,p \in \mathcal{U}} \|p - p_{\text{model}}(x; \theta)\|_2^2.$$

UNLABELED LOSS

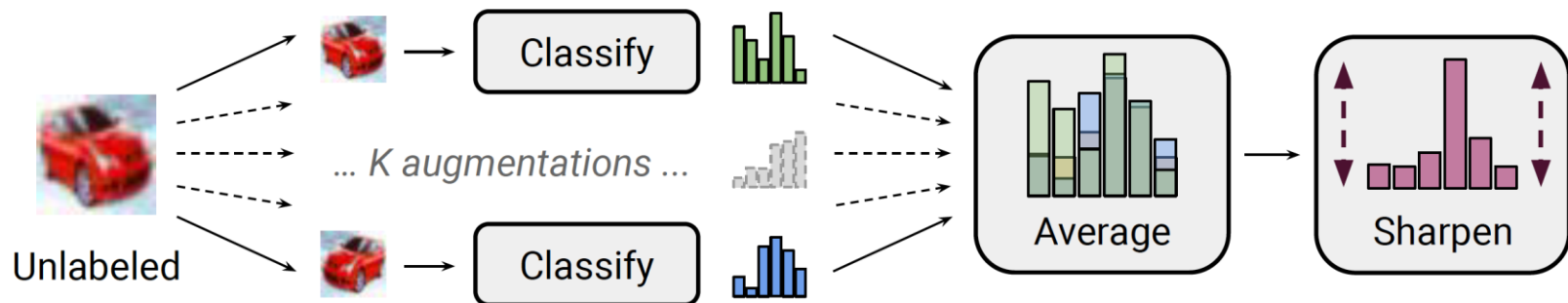
$$\mathcal{L}_{\text{reg}} = \sum_c \pi_c \log \left(\pi_c / \left(\frac{1}{|\mathcal{X}| + |\mathcal{U}|} \sum_{x \in \mathcal{X} \cup \mathcal{U}} p_{\text{model}}^c(x; \theta) \right) \right).$$

REGULARIZED LOSS

- Global noise modelling **assumes all samples as i.i.d**
- Noise leads to **some classes being harder to learn**
- Different classes exhibit **different loss characteristics**

Are the data enough?

The Mixup algorithm



Label guessing

$$\bar{q}_b = \frac{1}{K} \sum_{k=1}^K p_{\text{model}}(y \mid \hat{u}_{b,k}; \theta)$$

over K augmentations of u_b

Mixup

$$\lambda \sim \text{Beta}(\alpha, \alpha)$$

$$\lambda' = \max(\lambda, 1 - \lambda)$$

$$x' = \lambda' x_1 + (1 - \lambda') x_2$$

$$p' = \lambda' p_1 + (1 - \lambda') p_2$$

Uncertainty



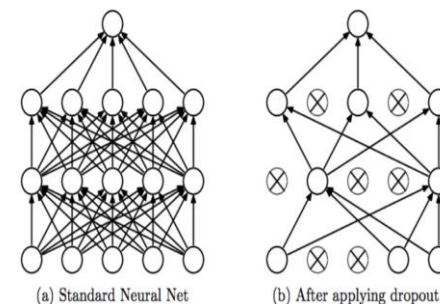
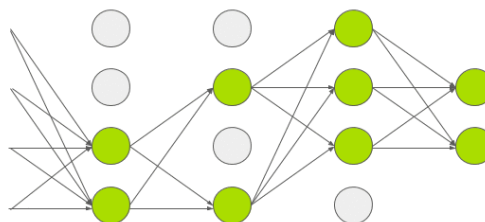
How certain is the model in making the predictions?



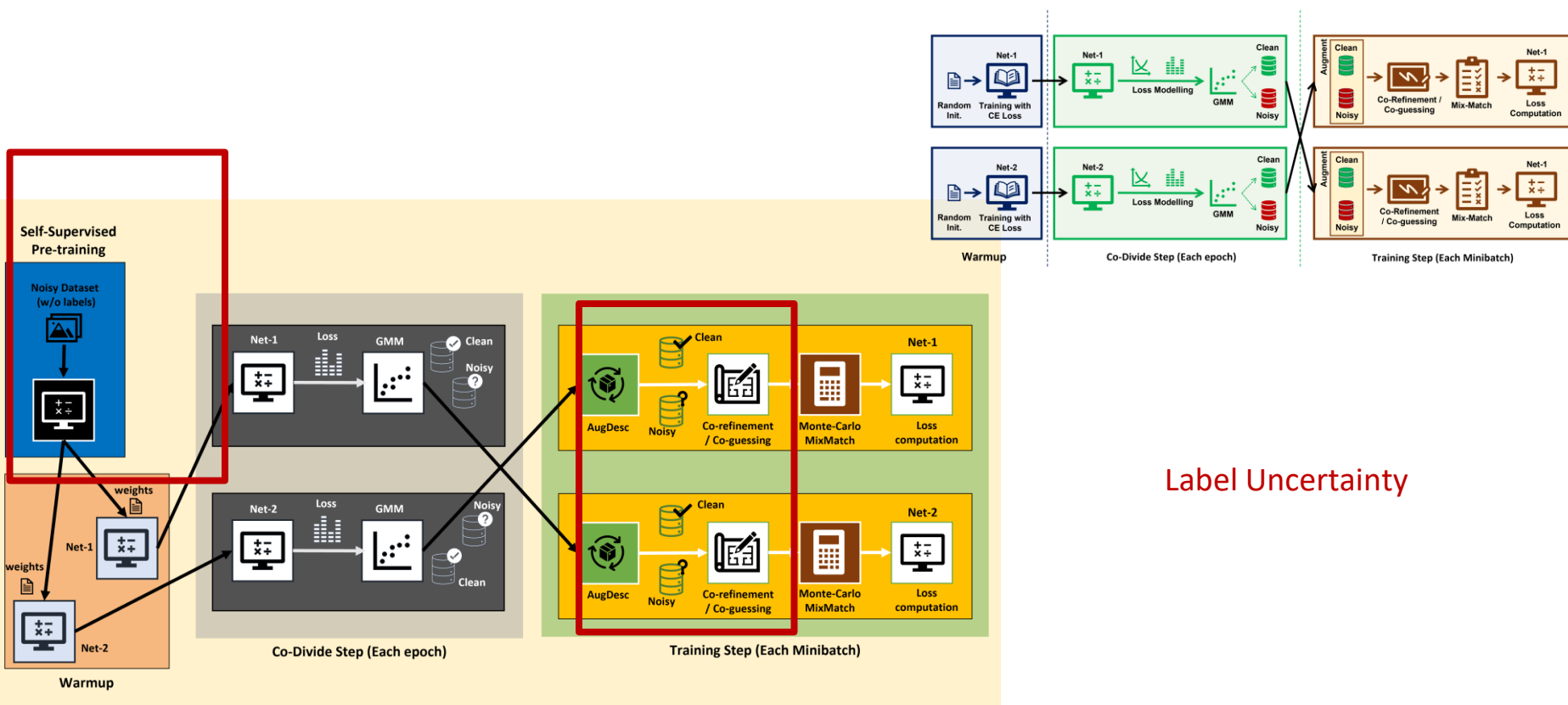
How to measure the uncertainty?

❖ Training a neural network with **dropout** is equivalent to doing approximate variational inference in a probabilistic deep Gaussian process.

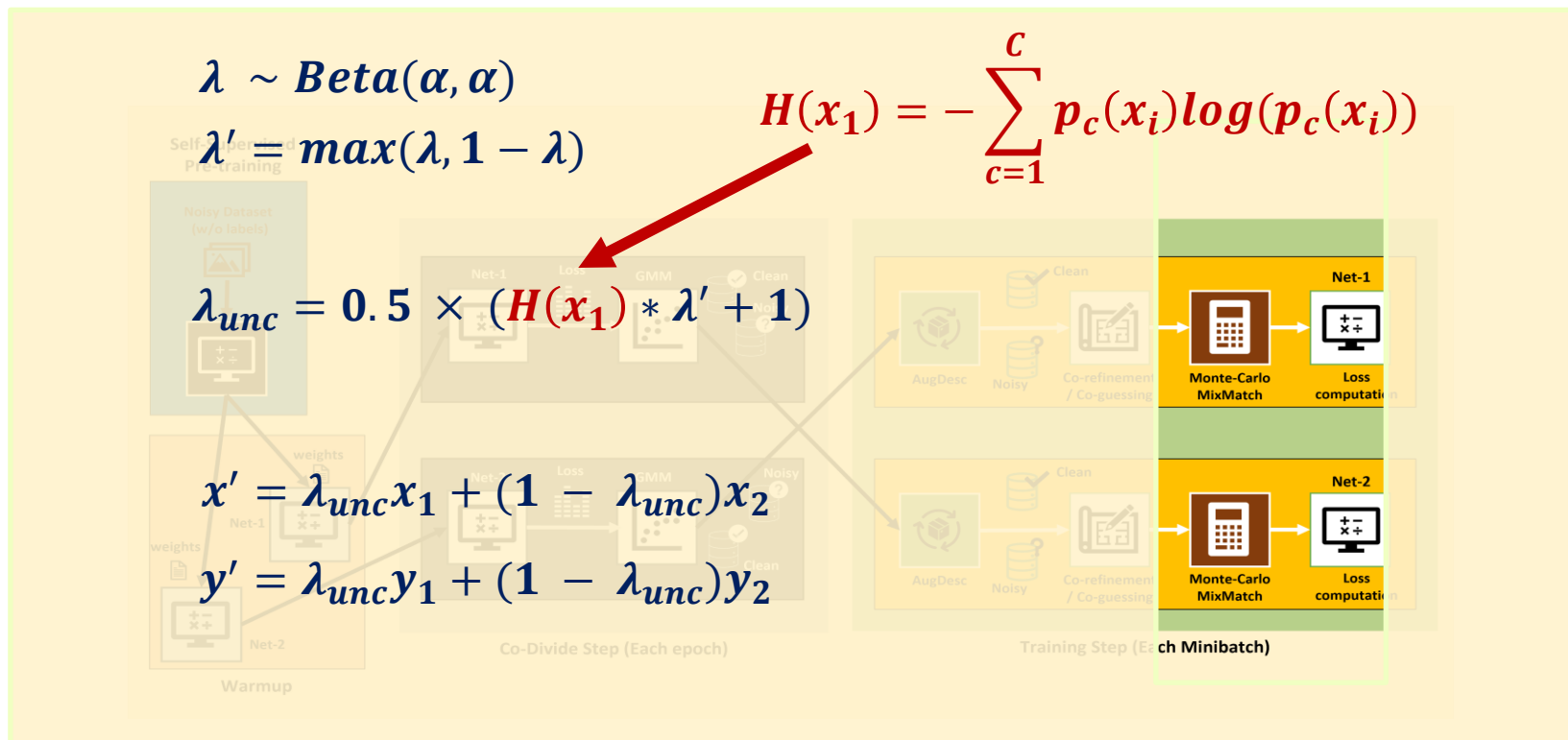
❖ When dealing with the predictive distribution, we can simply have a **Bernoulli distribution over the weights**.



Bayesian DivideMix++



Monte-Carlo MixMatch



Class-conditional Importance Weighted Loss

Cross Entropy Loss over
augmented and mixed clean data

Mean squared error over
augmented and mixed noisy data

$$\text{Total loss: } \mathcal{L} = \mathcal{L}_x + \lambda_u \mathcal{L}_u + \lambda_r \mathcal{L}_{reg}$$

Class-conditional
Supervised Loss

$$\mathcal{L}_x = - \frac{1}{|\mathcal{X}'|} \sum_{\mathbf{x}, \mathbf{p} \in \mathcal{X}'} \sum_c p_c \log(p_c^{\text{model}}(\mathbf{x}; \theta))$$

Regularization
term

Uncertainty-aware Loss

$$\mathcal{L}_u = \frac{1}{|\mathcal{U}'|} \sum_{(\mathbf{x}, \mathbf{p}) \in \mathcal{U}'} \|\mathbf{p} - \mathbf{p}_{\text{model}}(\mathbf{x}; \theta)\|_2^2$$

Unsupervised Loss

$$\mathcal{L}_{reg} = \sum_c \pi_c \log \left(\frac{\pi_c}{\left(\frac{1}{|\mathcal{X}'| + |\mathcal{U}'|} \sum_{\mathbf{x} \in \mathcal{X}' + \mathcal{U}'} p_{\text{model}}^c(\mathbf{x}; \theta) \right)} \right), \quad \pi_c = \frac{w_c}{C}$$

Class-conditional Loss

$$\mathbf{w} = \frac{\max(\lambda_w, f')}{|\max(\lambda_w, f')|}$$

Validation



Performance Comparison

Datasets	Method		20%	50%	80%	90%	Asym. 40%					
CIFAR-10	DivideMix	Best	96.1	94.6	93.2	76	93.4					
		Average	95.7	94.4	92.9	75.4	92.1					
	CCLM (Ours)	Best										
		Average	Method / Noise ratio		20%		50%		80%		90%	
CIFAR-100	DivideMix	Best	DivideMix (ICLR, 2020) Li et al.	Best	96.1		94.6		93.2		76.0	
				Last	95.7		94.4		92.9		75.4	
		Average	DM + AugDesc (CVPR, 2021) Nishi et al.	Best	96.3		95.4		93.8		91.9	
				Last	96.2		95.1		93.6		91.8	
	CCLM (Ours)	Best	DM + C2D (WACV, 2022) Zheltonozhskii et al.	Best	96.43±0.07		95.32±0.12		94.40±0.04		93.57±0.09	
				Last	96.23±0.09		95.15±0.16		94.30±0.12		93.42±0.09	
		Average	Sel-CL+ (CVPR, 2022) Li et al.	Best	95.5		93.9		89.2		81.9	
		UNICON (CVPR, 2022) Karim et al.	Best	96.0		95.6		93.9		90.8		
			Last	—		—		—		—		
		ULC (AAAI, 2022) Huang, Bai, et al.	Best	96.1		95.2		94.0		86.4		
Last			95.9		94.7		93.2		85.8			
LongReMix (PR, 2023) Cordeiro et al.		Best	96.3±0.1		95.1±0.1		93.8±0.2		79.9±2.7			
		Last	96.0±0.1		94.8±0.1		93.8±0.2		79.1±3.1			
Lipschitz Reg. (PR, 2023) Miao, et al.		Best	96.2		95.2		93.4		85.0			
		Last	95.7		94.8		93.1		84.3			
	Bayesian DivideMix++ (Ours)	Best	96.39±0.06		95.68±0.09		95.25±0.08		94.46±0.15			
		Last	96.13±0.07		95.40±0.11		94.97±0.02		94.20±0.12			

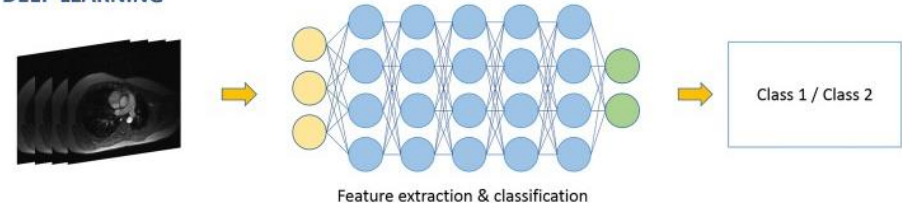
¿What makes NNs so popular?



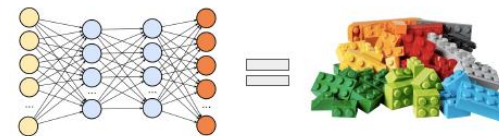
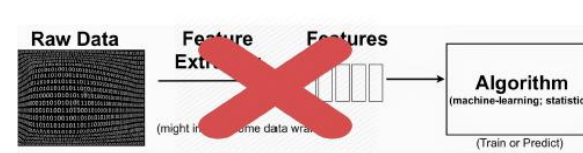
Advantages:

1. End-to-end learning

DEEP LEARNING



2. Modularity



3. Transfer learning



Imagine...

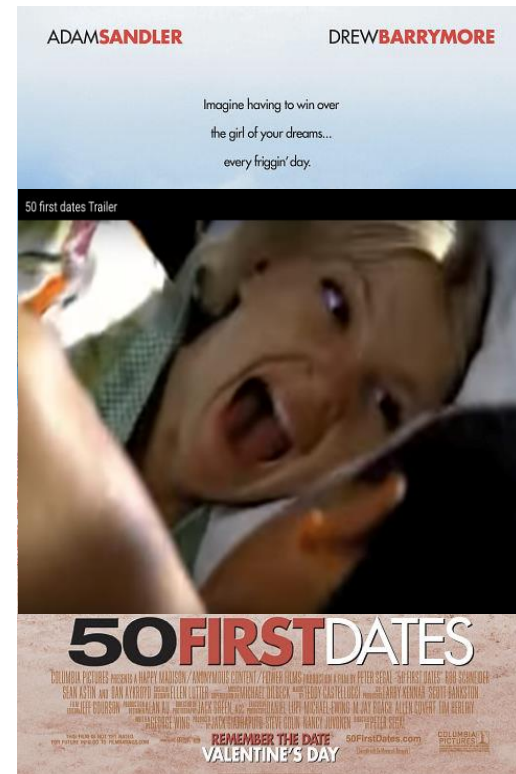


Henry Roth is a man afraid of commitment
until he meets the beautiful Lucy.

They hit it off, and Henry thinks he's finally
found the girl of his dreams.

Until he discovers:

she has short-term memory loss and forgets
him the next day.



Transfer Learning and Fine-tuning



Training data	1000s to millions of labeled images
Computation	Compute intensive
Training Time	Days to Weeks for real problems
Model accuracy	High (can over fit to small datasets)

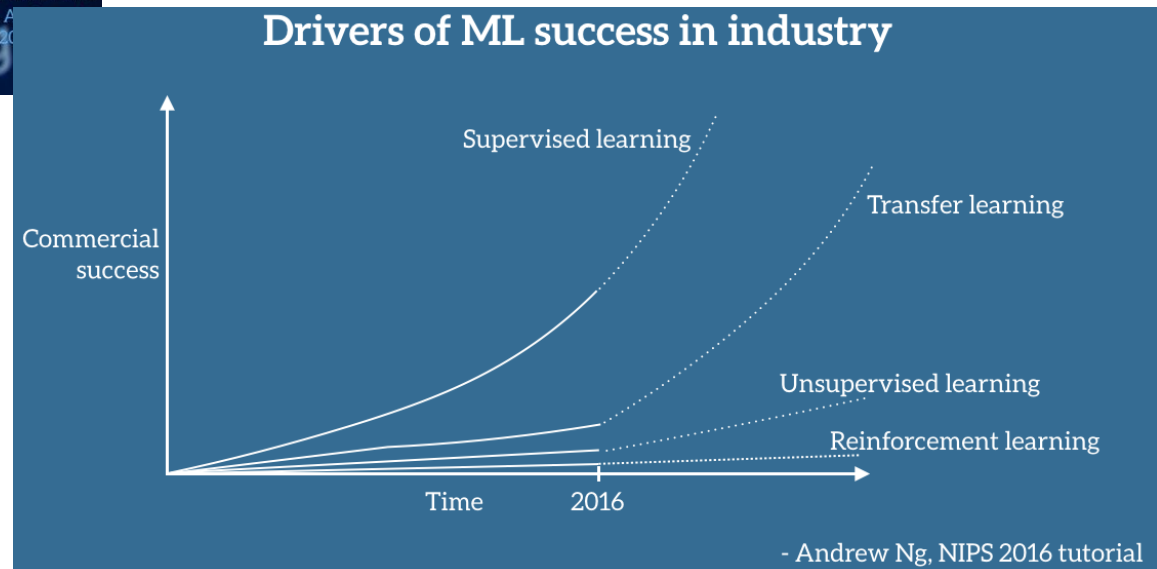
Training data	100s to 1000s of labeled images (small)
Computation	Moderate computation
Training Time	Seconds to minutes
Model accuracy	Good, depends on the pre-trained CNN model

Example of Fine-tuning

¿Why Transfer Learning?



Andrew Ng, chief scientist at Baidu and professor at Stanford



<https://ruder.io/transfer-learning/index.html#fn44>

The problem: A tale about machine learning

Once upon a time, a lot of data was collected.

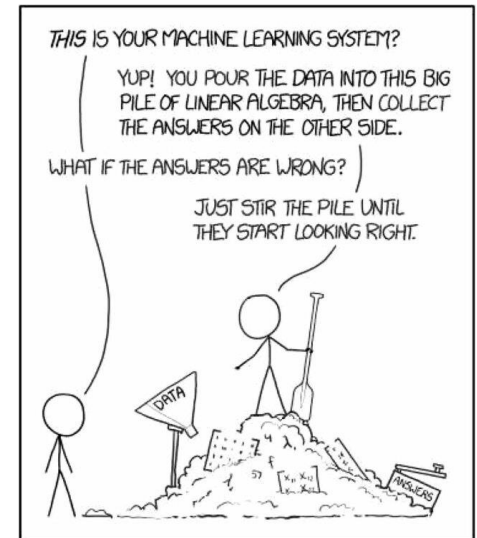
That data was fed into a huge machine learning model.

The model was then able to properly process the data, and when new similar data arrived at the model, it can correctly handle it.

What's wrong with this?

Nothing! It's just how machine learning works!

And it works wonderfully!



From: PhD Gabriele Grafietti

The problem: The antagonists of the tale



- **...a lot of data was collected is it always possible?**
 - How to store it?
 - What about training time?
 - What if I don't want to wait until I collected a lot of data?
 - When some data is a lot? When some data is sufficient?
- **...when new similar data arrived at the model, it can correctly handle it...**
 - What if dissimilar data arrives at the model?
 - What if I want to adapt the model to new data?
 - What if data changes over time?
 - What if I want to continually train the model?

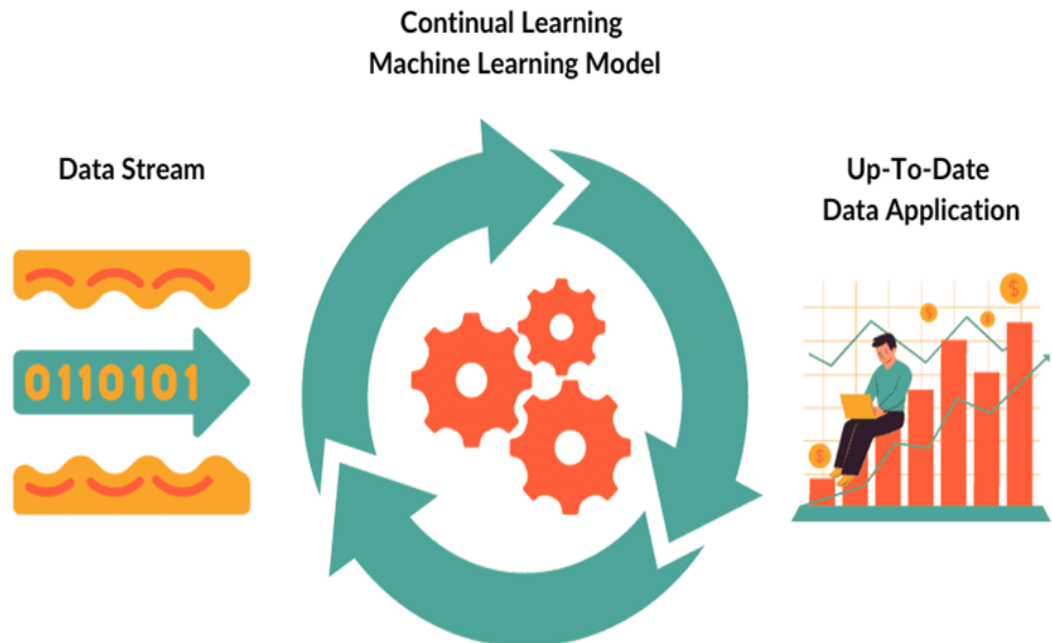
Continual learning



Continual Learning aims to enable the model to learn continuously from new data streams, making the **process more efficient and scalable**.

Training from scratch, every time new data arrives, is costly

Privacy concerns may make old data inaccessible



The main villain of ML: catastrophic forgetting

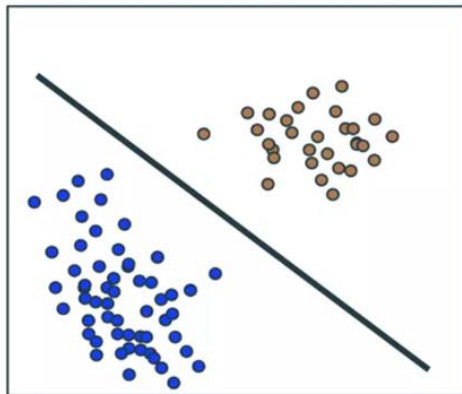


- **Catastrophic interference**, also known as **catastrophic forgetting**,
 - is the tendency of an NN to completely and abruptly forget previously learned information upon learning new information.
 - This holds for all the ML models that are trained using "greedy" algorithms (e.g. SGD, CART...)
- The training algorithm **optimizes** the parameters of the model using the **currently available** data,
 - **past data and past knowledge** are not taken into consideration.

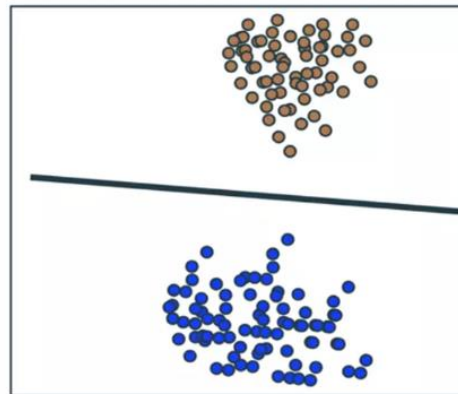
Catastrophic Forgetting: an example



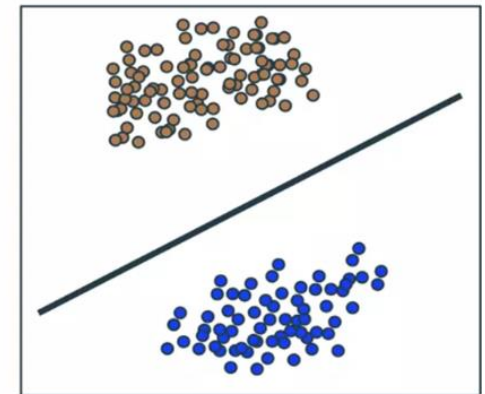
- You collect some data from the real world, and after each collection phase you train the model (only on the lastly collected data) to classify it in two classes:



Data 1



Data 2

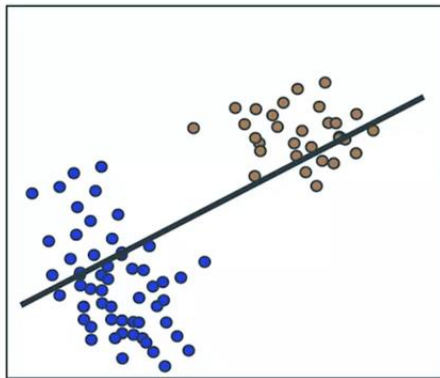


Data 3

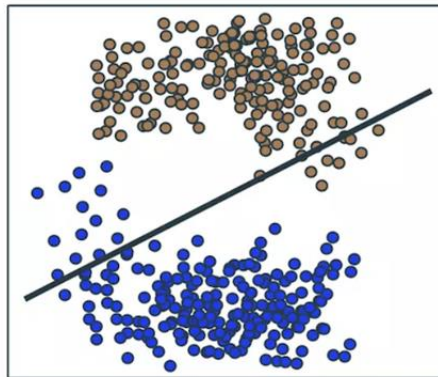
Catastrophic Forgetting: an example



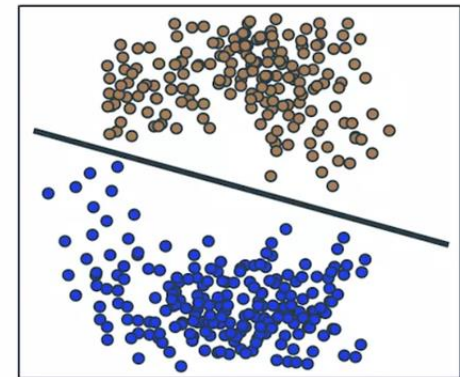
- You collect some data from the real world, and after each collection phase you train the model (only on the lastly collected data) to classify it in two classes:



Final solution on data 1



Final solution on all data



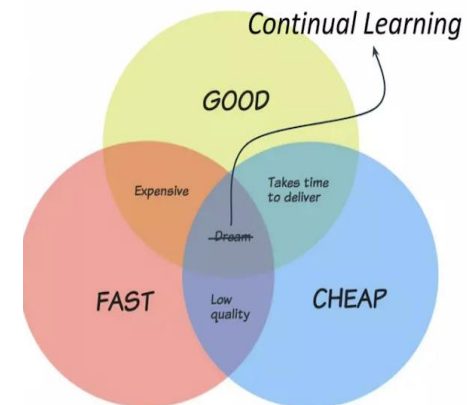
Optimal final solution

Continual learning



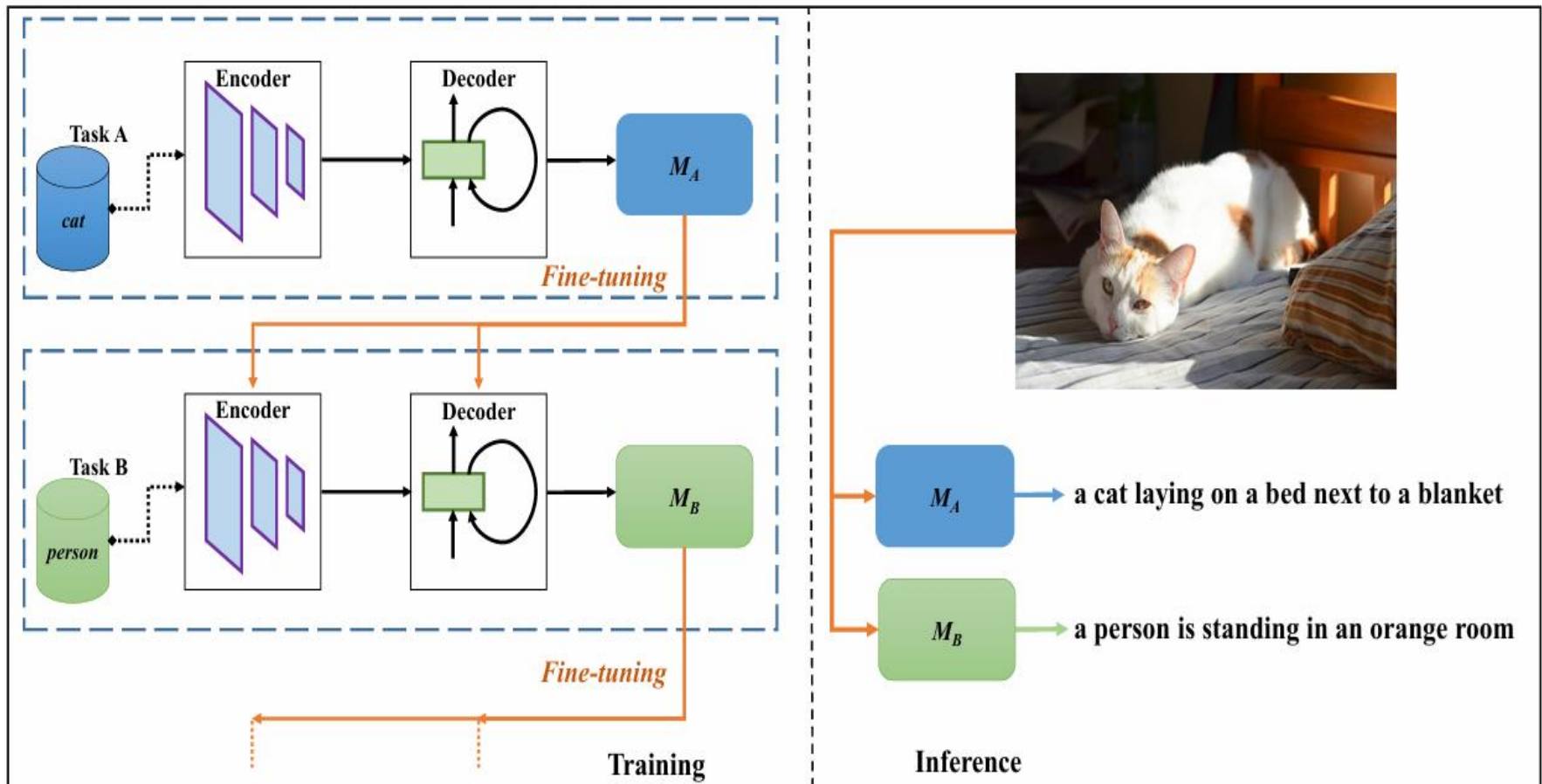
- **Continual learning** - a scenario when we do not have all the data at once, but we discover new data as time progresses.
 - The data we discover may not be a good approximation for the total data distribution.

- Other constraints:
 - Every time new data arrives the **model needs to be updated**
 - The model update **should be fast** enough to be used before new data arrives
 - **Past knowledge must not be forgotten** (at least not catastrophically)



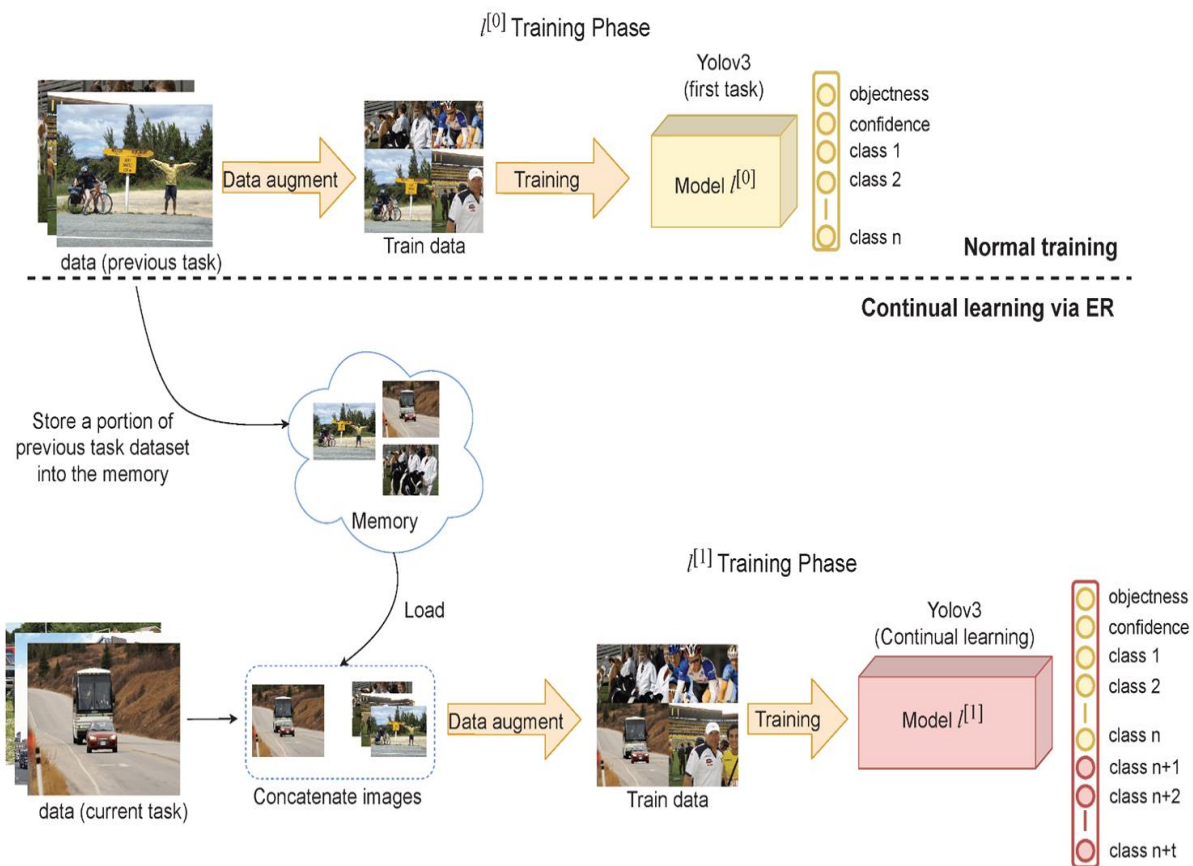
From: PhD Gabriele Graftetti

Catastrophic Forgetting in Continual Learning



Replay-based as an effective strategy to retain prior knowledge

Replay-based CL approaches retain previous knowledge by maintaining and revisiting a small buffer.



Replay strategy

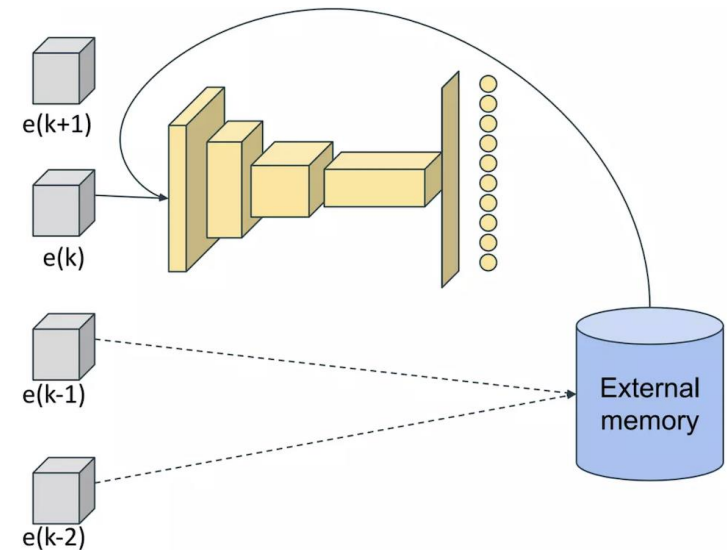


Pro:

- Catastrophic forgetting highly reduced.
- Simple and easy to implement strategy.
- Memory is cheap and abundant.

Cons:

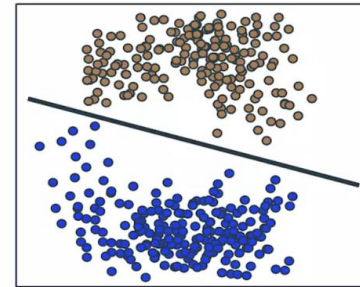
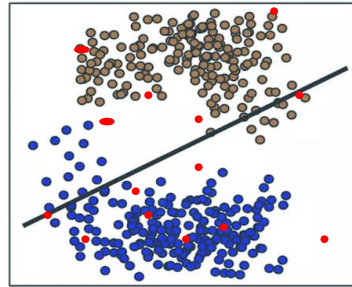
- External memory
- Memory is not infinite (the stream of experience can be infinite).
- What about privacy and private data?
- Not biologically plausible.
- Computation



From: PhD Gabriele Graftetti

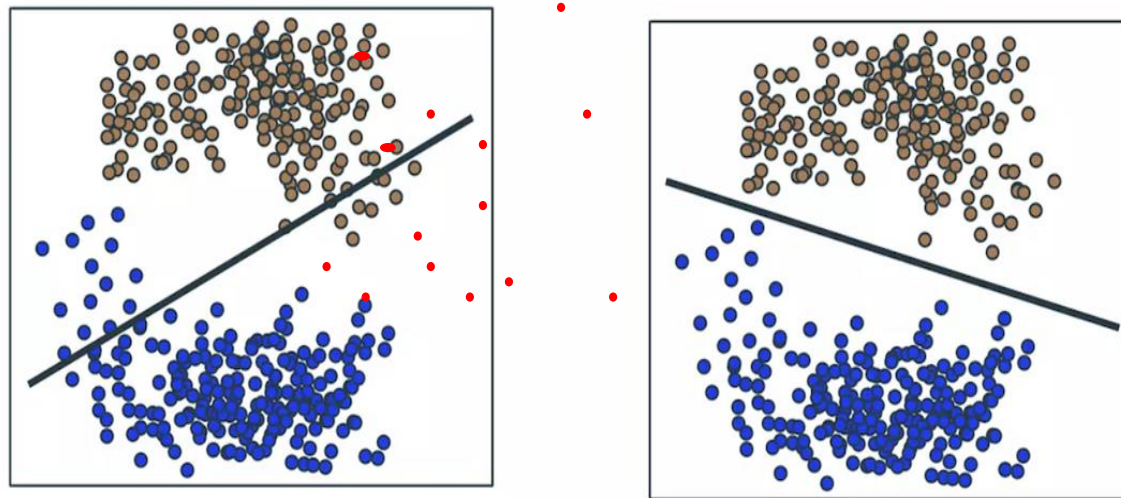


But.... which data to retain?

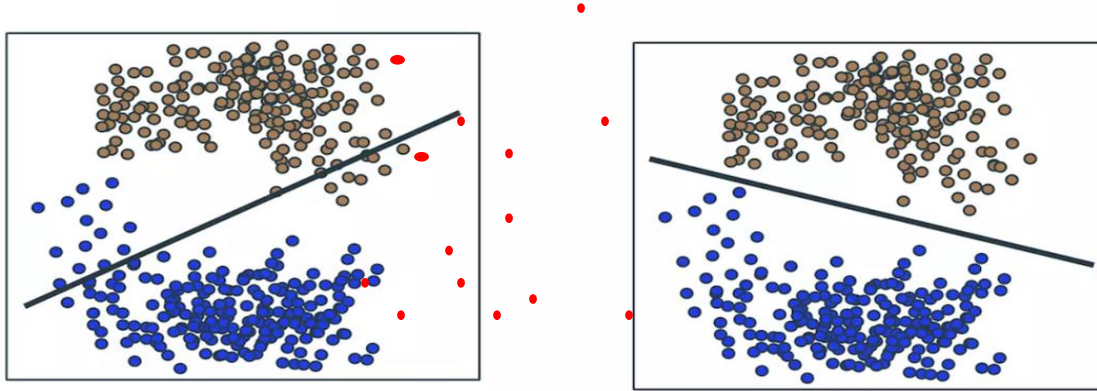


The simplest method **Random Selection**, continuous to be the one commonly chosen in CL

- **Still? Completely random?**
- Uncertainty-based approaches have proven very effective in improving the understanding of the DL models
- Data with high epistemic uncertainty means being underrepresented



Our hypothesis: the uncertainty score related to each sample may be a good indicator when selecting a suitable example to retain prior knowledge.



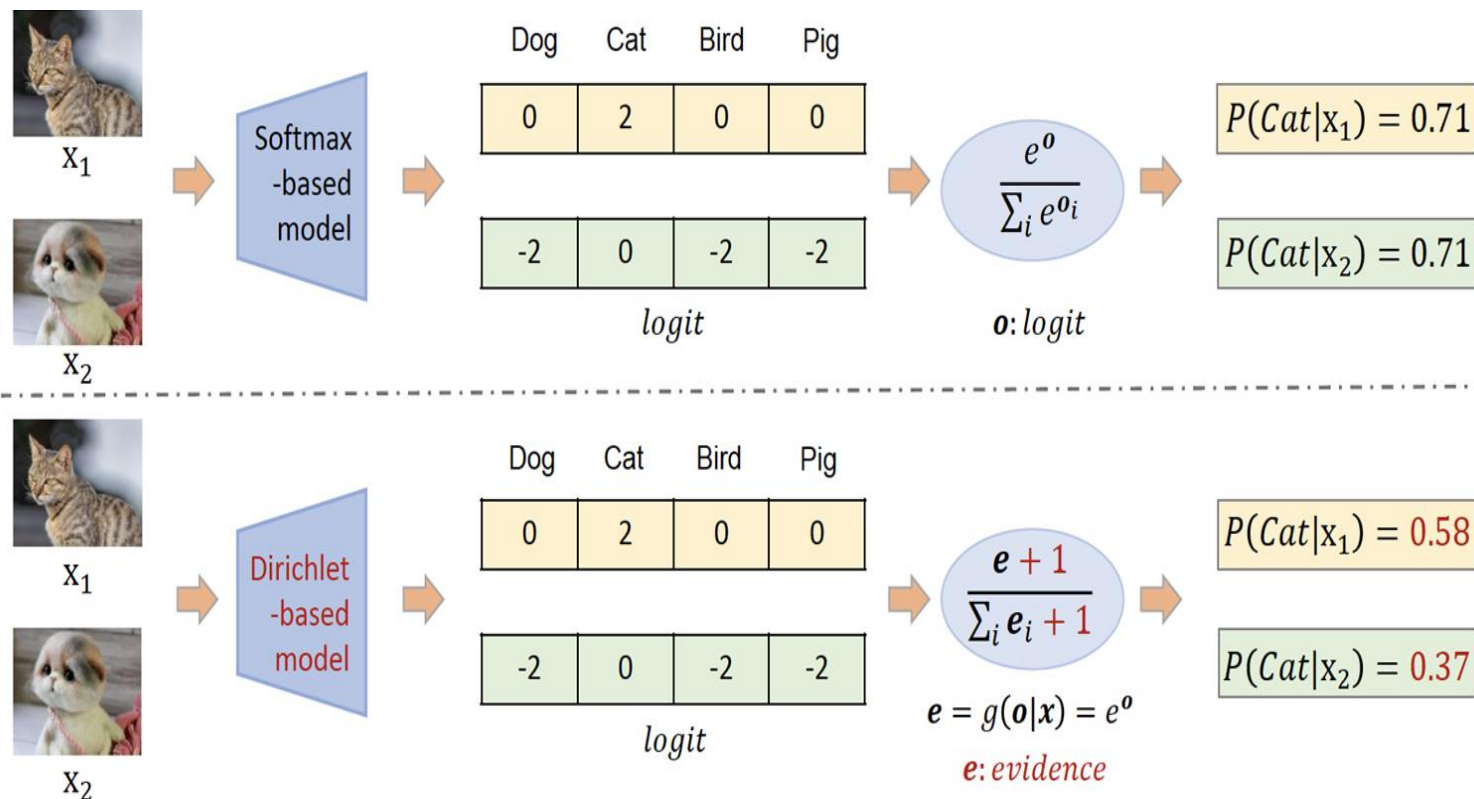
But how to compute the Uncertainty of a predictive model?

- **Evidential Deep Learning method**

Evidential Deep Learning Framework

Let's a predictive model $f_\theta(x_i)$ give the evidence or prediction e_i for a sample x_i , σ is $\exp()$.

$$e_i = \sigma(f_\theta(x_i))$$



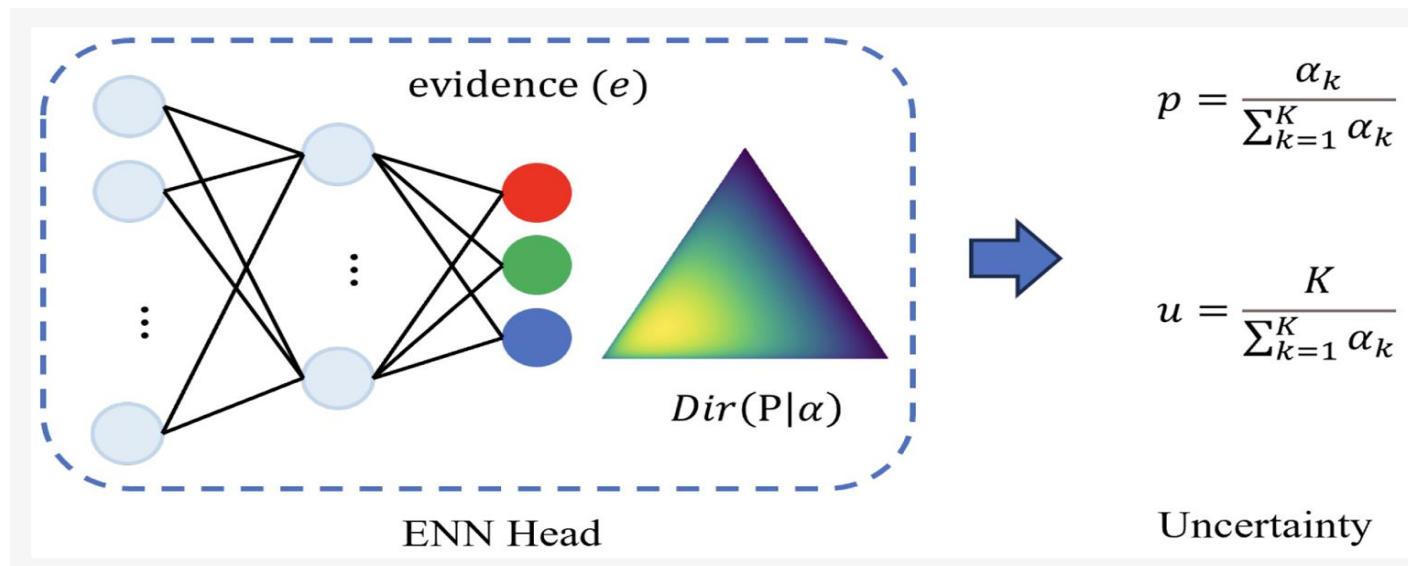


Parameters of a Dirichlet distribution:

$$\alpha_i = e_i + 1.$$

Dirichlet strength (belief):

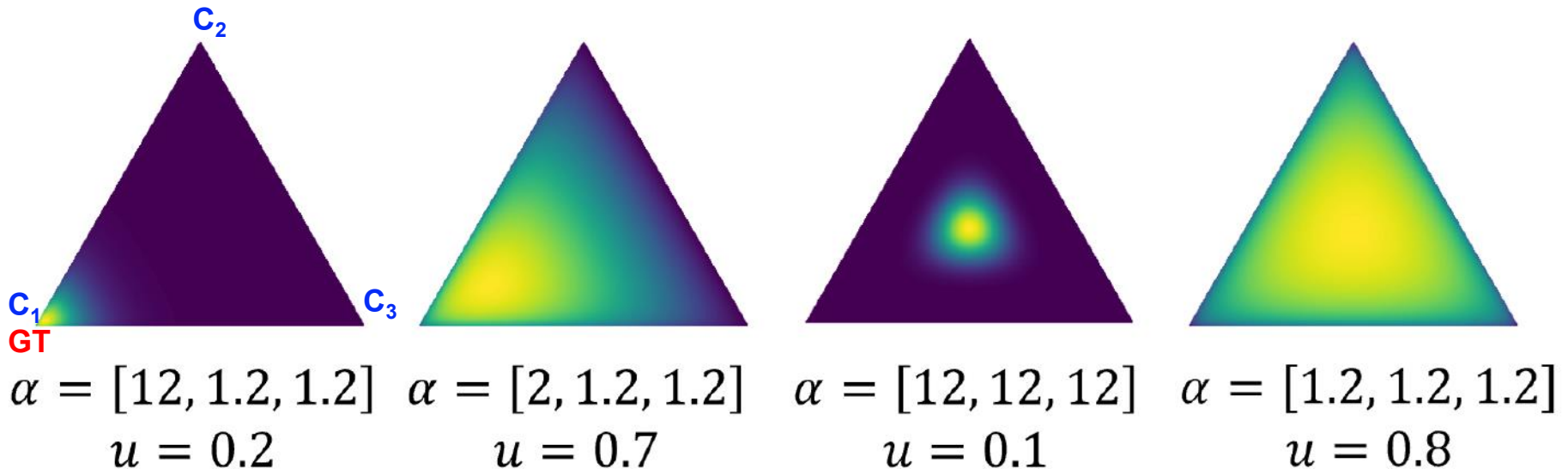
$$S_i = \sum_{j=1}^K \alpha_{ij} \quad K \text{ is the number of classes.}$$



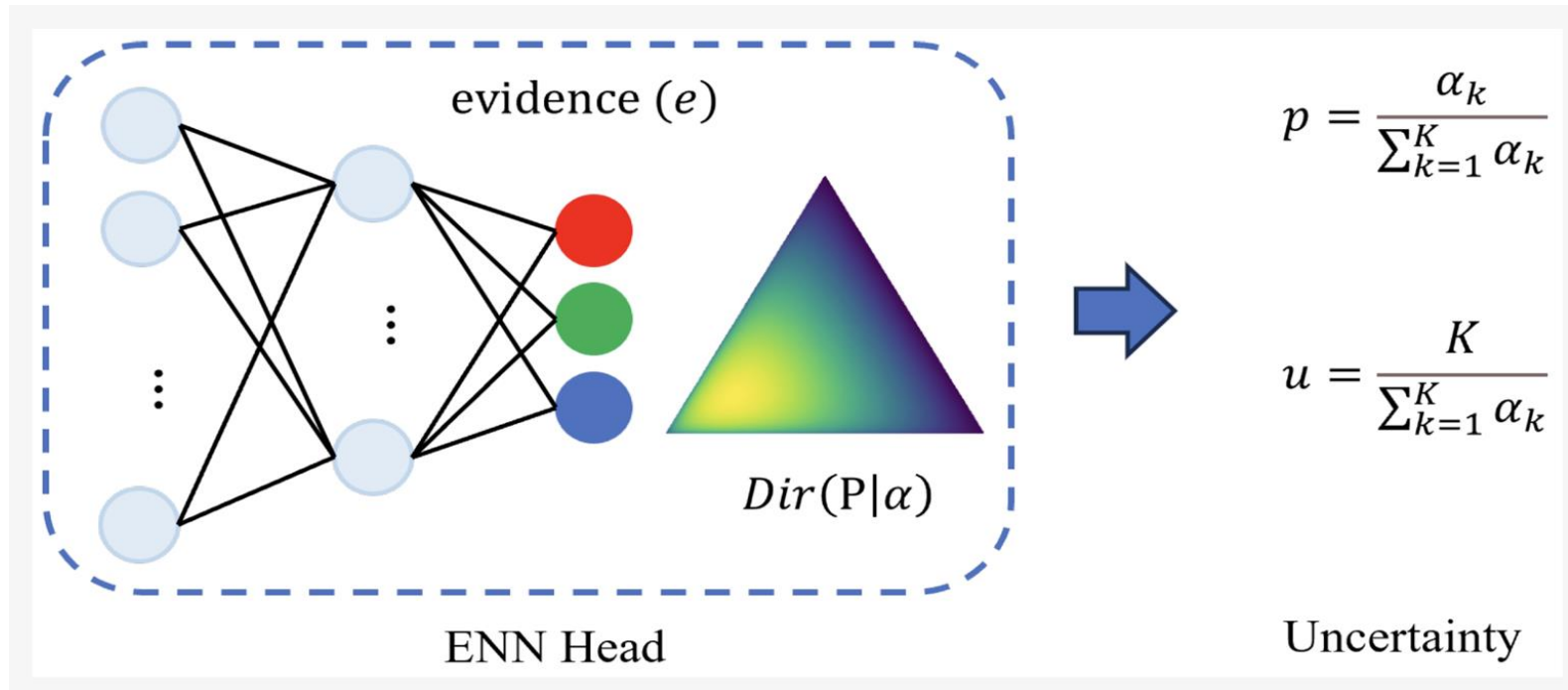


Uncertainty:

$$u = \frac{K}{\sum_{k=1}^K \alpha_k}$$



Accurate & certain, Accurate & uncertain, Inaccurate & certain, Inaccurate & uncertain



Dirichlet strength (S_i):

$$S_i = \sum_{j=1}^K \alpha_{ij}$$

Evidential DL Loss:

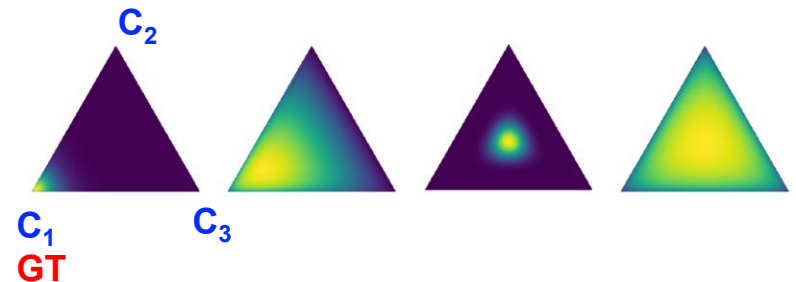
$$L_i = \sum_{j=1}^K y_{ij} \times (\log(S_i) - \log(\alpha_{ij}))$$



Regularization term controls the evidence of misclassified samples:

$$\alpha_{KL_{ij}} = e_{ij} \times (1 - y_{ij}) + 1; \quad S_{KL_i} = \sum_{j=1}^K \alpha_{KL_{ij}}$$

$$KL_{1_i} = \ln \frac{\Gamma(S_{KL_i})}{\Gamma(K)} - \sum_{j=1}^K \ln \frac{\Gamma(\alpha_{KL_{ij}})}{\Gamma(1)}$$



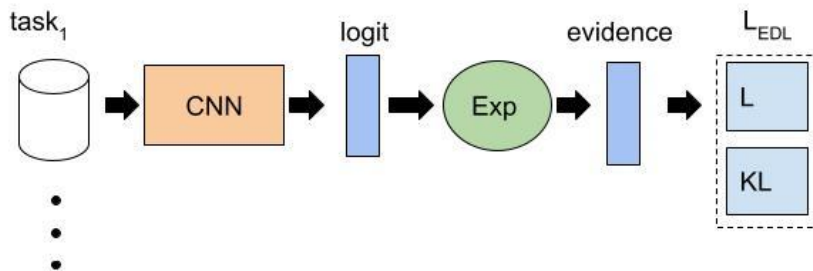
$$\Gamma(n) = (n - 1)!$$

$$KL_{2_i} = \sum_{j=1}^K (\alpha_{KL_{ij}} - 1) \times \left(\frac{\Gamma'(\alpha_{KL_{ij}})}{\Gamma(\alpha_{KL_{ij}})} - \frac{\Gamma'(S_{KL_i})}{\Gamma(S_{KL_i})} \right)$$

$$KL_i = KL_{1_i} + KL_{2_i}.$$

<https://statproofbook.github.io/P/dir-kl.html>

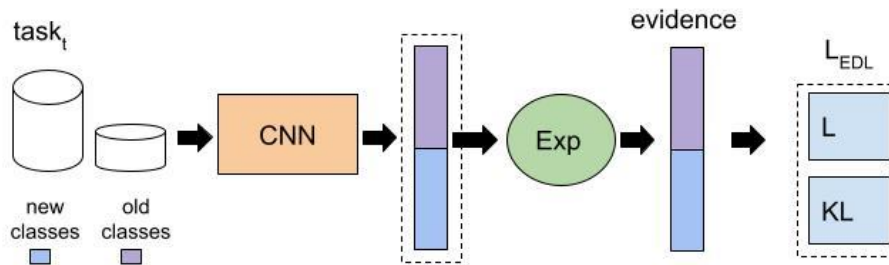
Continual Evidential Deep Learning Framework



Evidential DL loss



$$L_i = \sum_{j=1}^K y_{ij} \times (\log(S_i) - \log(\alpha_{ij}))$$



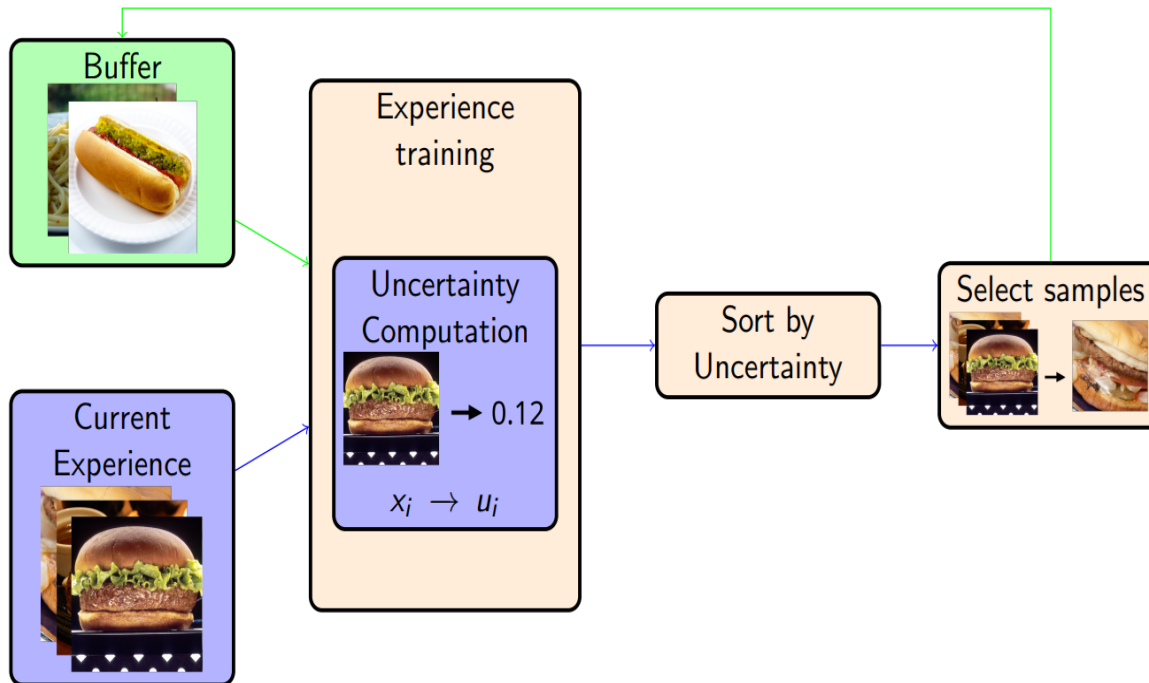
Missclassified KL loss

$$KL_i = KL_{1_i} + KL_{2_i}.$$

Total EDL loss:

$$L_{EDL_i} = (1 - \lambda) \times L_i + \lambda \times (C_{ann} \times KL_i).$$

Uncertainty-based Selection Methodology

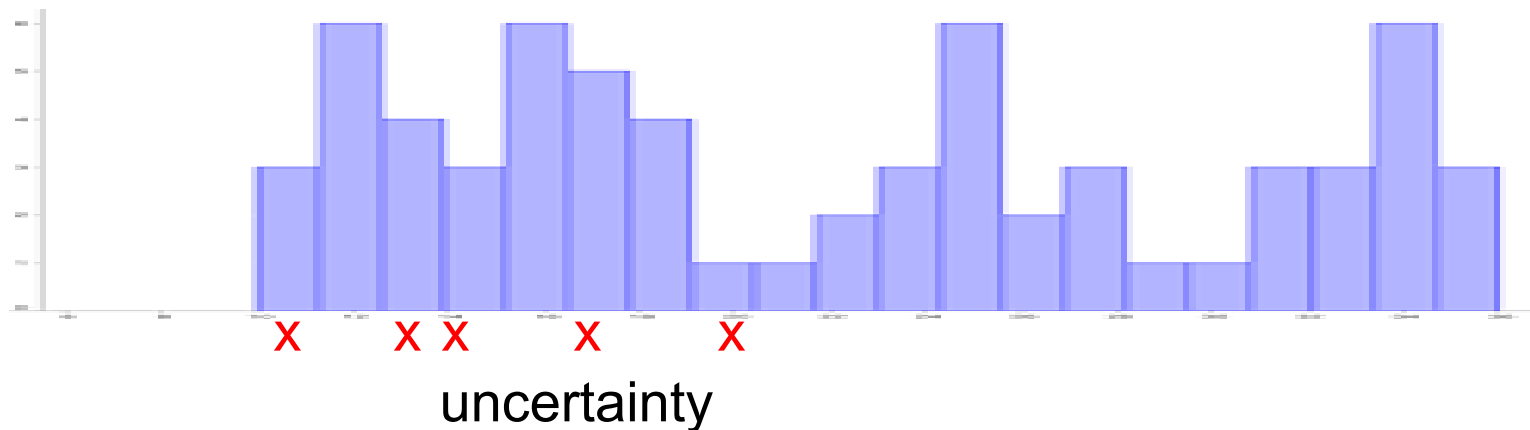


EDL-based Rehearsal methods

- Samples Sorting
- Samples Clustering
- Samples Filtering

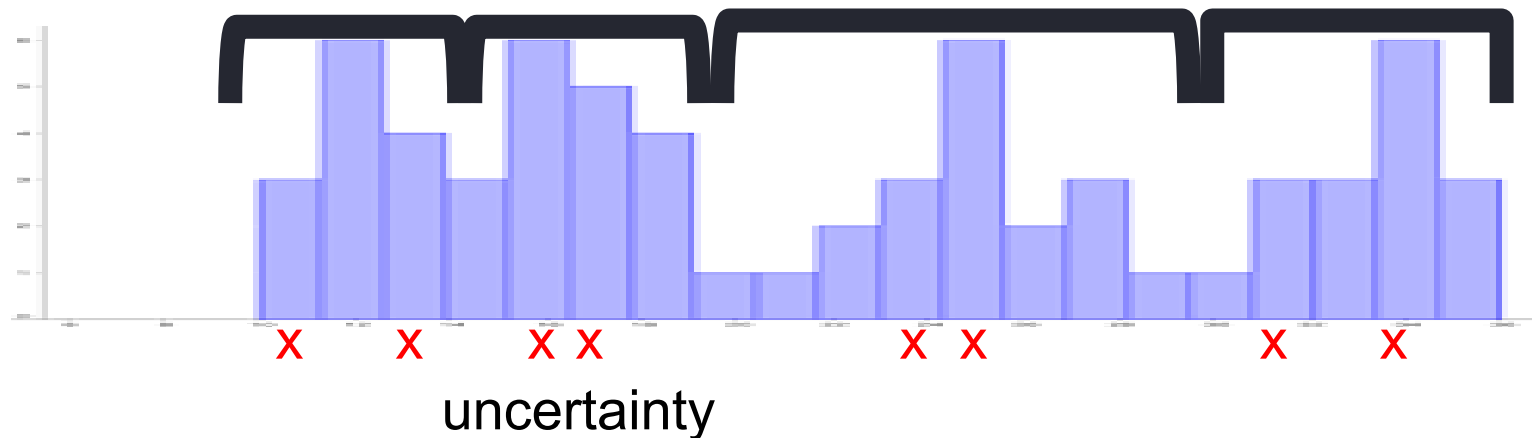
Methodology - EDL-based Rehearsal methods

- **Samples Sorting:** Sort the samples according to their uncertainty and select those with the lowest uncertainty.



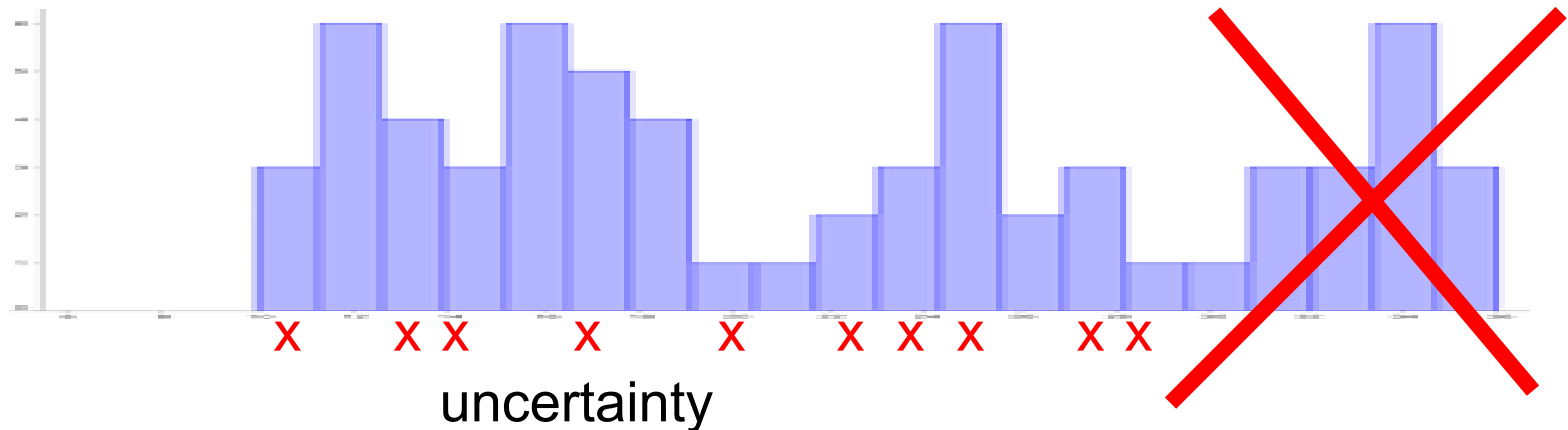
Methodology - EDL-based Rehearsal methods

Samples Clustering: apply K-Means, select randomly (*Kmeans random*) within each cluster or according to their proximity to the median (*Kmeans median*).



Methodology - EDL-based Rehearsal methods

- **Samples Filtering:** eliminate the samples with the highest uncertainty and choose the samples randomly (*Filtered random*) or by means *Kmeans random* (*Filtered Kmeans*).



Experimental design



Datasets:

- CIFAR10
- CIFAR100
- FOOD101

Model:

- ResNet-18

Training parameters:

Dataset	Epochs	Increments	Base
<i>CIFAR10</i>	10	2	2
<i>CIFAR100</i>	50	20	20
<i>FOOD101</i>	120	20	21

airplane

automobile

bird

cat

deer

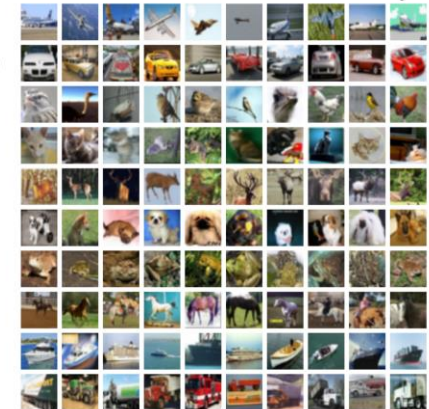
dog

frog

horse

ship

truck



Experiments - Evaluation Metrics



- **Final Accuracy**

$$Acc_i^j = \frac{TP_i^j + TN_i^j}{TP_i^j + TN_i^j + FP_i^j + FN_i^j}$$

- **First Experience Classes Final accuracy**

$$AccFinal = Acc_{1,...,nE}^{nE}; \quad Acc1st = Acc_1^{nE}$$

- **Forgetting**

Example:

T1 T2 T3 T4

1 step: 80

2 step: 75 85

3 step: 72 80 90

4 step: 70 80 80 85

$Acc_{final} = (70+80+80+85)/4=78,75$

$Acc1st = (80+85+90+85)/4=85$

$Forgetting = Acc1st - Acc_{final} = 85 - 78,75 = 6,25$

$$Forgetting = \frac{1}{nE} \sum_{i=1}^{nE} Acc_{exp_i} - Acc_{final_i}$$

Results - Baseline comparison

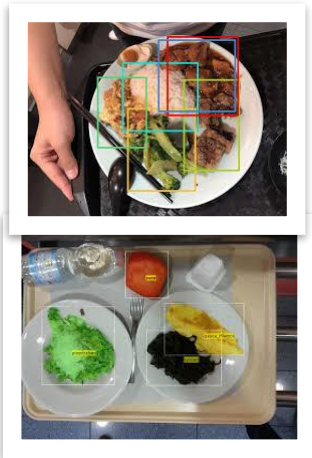


Dataset	Strategy	$AccFinal \uparrow$	$Acc1st \uparrow$	$Forgetting \downarrow$
CIFAR10	Random EDL	0.8927 ± 0.0275	0.8575 ± 0.0910	0.0571
	Filtered Random	0.8975 ± 0.0113	0.8653 ± 0.0842	0.0531
CIFAR100	Random EDL	0.5544 ± 0.0112	0.4824 ± 0.0491	0.2856
	Filtered Random	0.5670 ± 0.0127	0.5241 ± 0.0374	0.2739
FOOD101	Random EDL	0.4855 ± 0.0353	0.4491 ± 0.0261	0.3311
	Filtered Random	0.5076 ± 0.0164	0.4795 ± 0.0407	0.2898

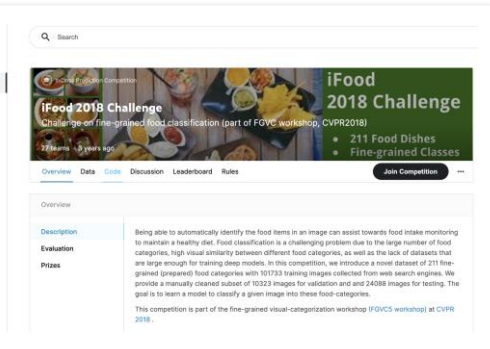
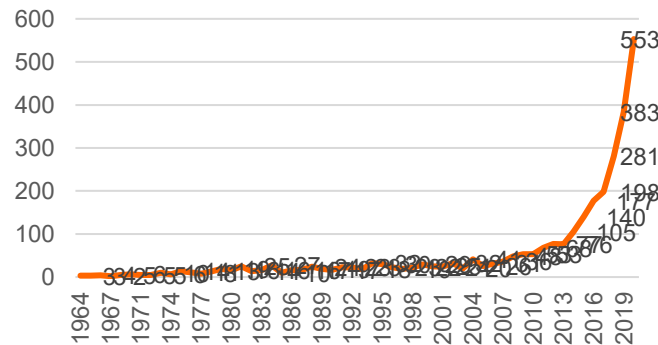
What application do we consider?



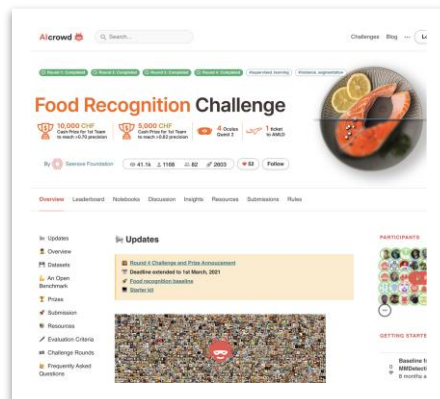
Food recognition popularity



Number of Food recognition papers



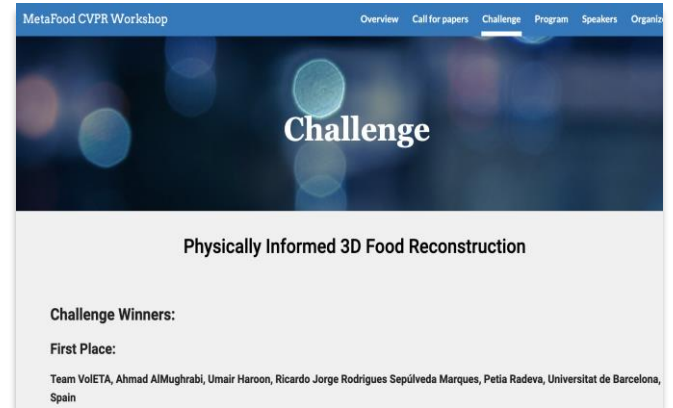
iFood 2011 fine-grained (prepared) food categories with 135733



AICrowd: 26000 annotated segmented images



LargeFineFoodAI: 1,000 fine-grained food categories and over 50,000 images.



MetaFood CVPR Workshop, 20 scenes, 5K images

Data-centric Food image analysis

Uncertainty modelling

Fine-grained Food recognition

Large-scale Food recognition

Food image analysis

Self-supervised Learning

Learning with Noisy Labeling Food recognition

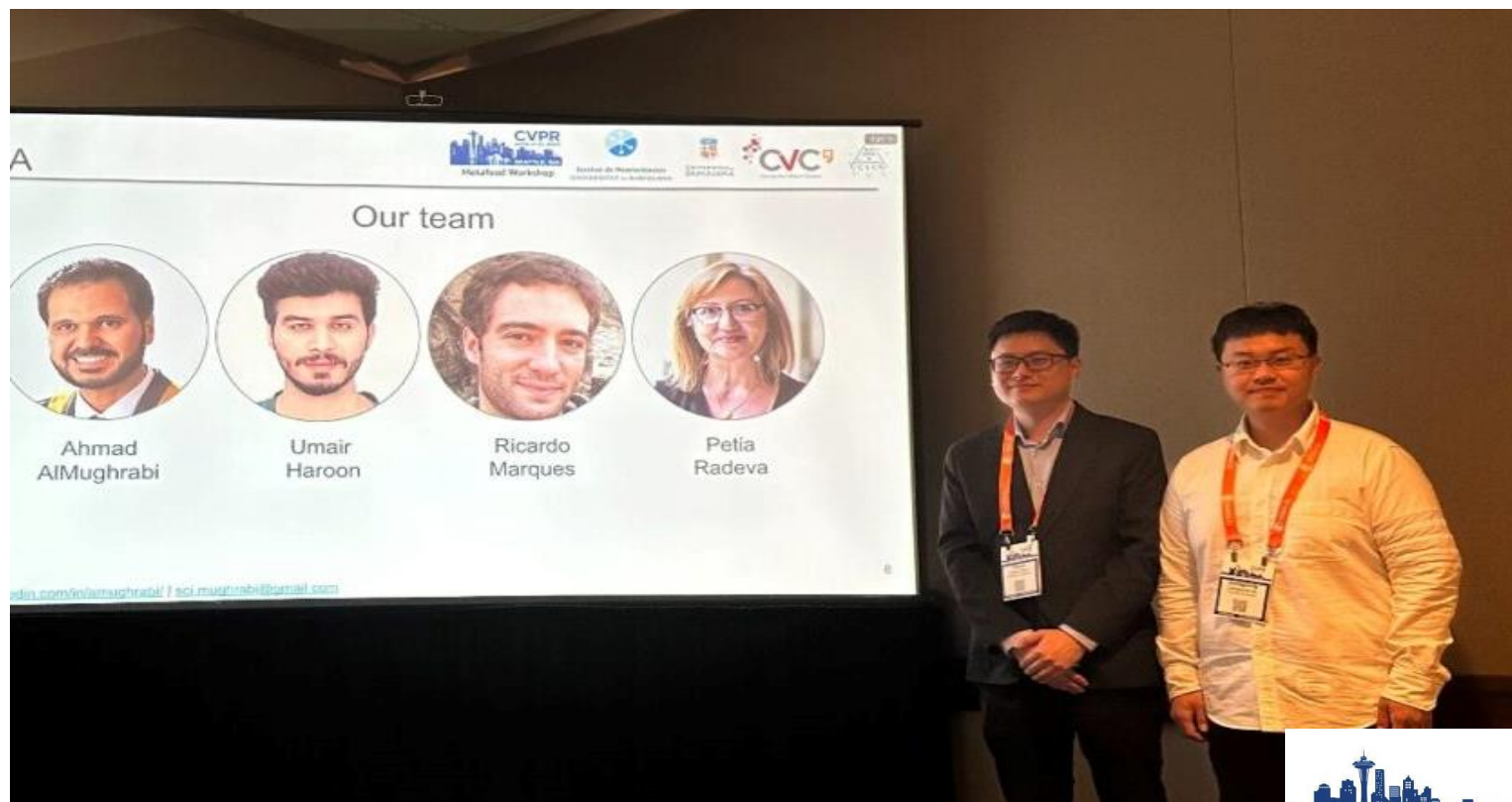
Generative AI for Food Volume Estimation

Food ontology-based Deep learning

Food Volume Estimation



CVPR MetaFood Workshop Challenge



The challenge aimed to **advance the food vision community by focusing on 3D food analysis, specifically nutrition intake monitoring through volume estimation**. The winners were revealed on June 17th, 2024



Workshop Challenge
WINNER



Success cases: European projects



- **Validity:** FOOD intake monitoring of kidney transplant patients



- **Nestore:** FOOD intake monitoring for malnutrition prevention in elderly



- **Greenhabit:** Greenhabit: a serious game to promote change behaviour

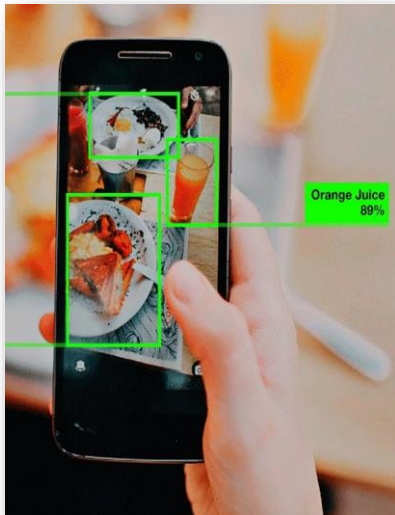


- **Diacare:** mHealth app to assist diabetic patients



Success story: Aigecko Technologies

API and APP that allows food recognition (ready meals and food) using Artificial Intelligence algorithms from just a photo




UNIVERSITAT DE BARCELONA

Noticias
Agenda
UBtv
Sal

APRENDE
INVESTIGA



Noticias

Inicio > Noticias > Nace AIgecko Technologies, la inteligencia artificial al servicio del reconocimiento...

Nace AIgecko Technologies, la inteligencia artificial al servicio del reconocimiento de imágenes



20/01/2021
Recerca

Ha nacido la nueva *spin-off* AIgecko, que ofrece servicios de reconocimiento y análisis de imágenes basados en los algoritmos desarrollados por miembros del Grupo de Investigación Computer Vision and Machine Learning de la UB, y que han dado pie a diferentes productos. La actividad de la empresa se enmarca en el aprendizaje profundo (*deep learning*), un campo de la inteligencia artificial que en los últimos años ha revolucionado el mercado y nuestro día a día en aspectos como la conducción automática, la visión artificial o los asistentes virtuales, entre muchas otras aplicaciones.

Más información

De izquierda a derecha, Marc Bolaños, Petia Radeva, M.ª Carme Verdaguer y Eric Verdaguer.

Conclusions

- ❖ **Learning with Noisy Labelling** based on Prompt Learning and class-conditional estimation improves the SoA
 - ❖ Division of clean-noisy samples central to any LNL sample selection algorithm
 - ❖ **CCLM uses local class division** of clean and noisy samples
 - ❖ **Bayesian DivideMix** proposes an elegant uncertainty-aware framework into the decision-making process
- ❖ Food recognition is a perfect test domain for powerful Machine/Deep Learning models
- ❖ **Food** Image Analysis is **highly underexplored problem** that could convert in an important **benchmark** for CV algorithms.
- ❖ **Continual learning** – an important field to assure robust transfer learning without forgetting
- ❖ Multiple other CV problems are highly relevant for food image analysis
- ❖ Multiple **real applications** and **professional opportunities**

1. Bhalaji Nagarajan*, Ricardo Marques*, Marcos Mejia, and Petia Radeva. "Class-conditional Importance Weighting for Deep Learning with Noisy Labels." VISAPP, 2022.
2. Albert Tatjer*, Bhalaji Nagarajan*, Ricardo Marques, and Petia Radeva. "CCLM: Classconditional Label Noise Modelling." IbPRIA, 2023, Best paper award.
3. Albert Tatjer*, Bhalaji Nagarajan*, Ricardo Marques, and Petia Radeva. "Decoding Class Dynamics in Learning with Noisy Labels." Pattern Recognition Letters, 2024..
4. Bhalaji Nagarajan, Ricardo Marques, Eduardo Aguilar, and Petia Radeva. "Bayesian DivideMix++ for Enhanced Learning with Noisy Labels." Neural Networks, 2024.
5. Aguilar, E., Raducanu, B., Radeva, P., & van de Weijer, J. (2025). CEDL+: Exploiting evidential deep learning for continual out-of-distribution detection. *Expert Systems with Applications*, 127774.

Is everything in DL done?

1. **Multimodal Learning:** Combining different sources (text, images) to create comprehensive models for understanding and generating content.
 - LLM and VLM for Dataset generation
 - Ontology-driven Deep Learning
2. **Transfer Learning:** these techniques became more widespread, as they can significantly reduce the amount of data required to train models.
 - Multi-task learning, Continual learning
3. **Uncertainty modeling:** refers to the process of quantifying and managing uncertainty or ambiguity in the predictions or decisions made by ML models
 - Uncertainty-aware MTL, Learning with noisy labeling
4. **Self-Supervised Learning:** methods train models on unlabeled data, which can be particularly useful in cases where labeled data is scarce.



Trends in DL

5. Efficiency and Model Compression: growing interest in making DL models smaller, faster, and more energy-efficient.

- Efficient learning based on low-rank models
- Scaling by hierarchical DL models
- Fine-grained recognition

6. Generative AI: GANs are used extensively, from image generation to data augmentation and domain adaptation.

- Uncertainty-aware data augmentation, NeRFs, Adversarial attacks

7. Explainable AI (XAI): The need for understanding and interpreting deep learning models became more critical, especially in fields like healthcare.

- Robust Explainable models

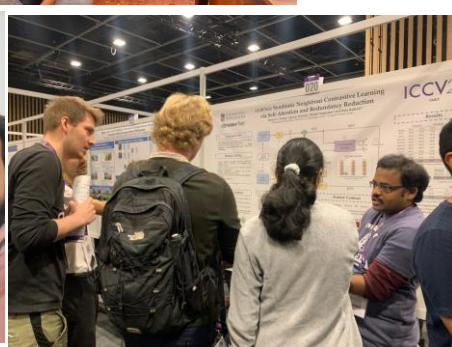
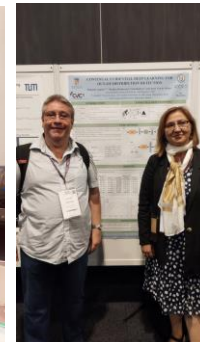
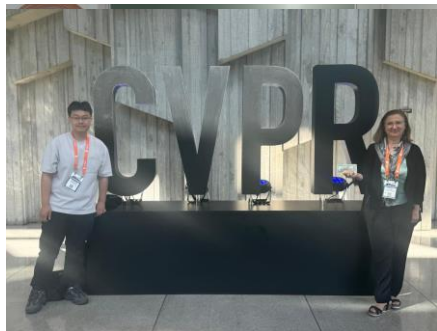
8. AI Applications (Healthcare, Transport, Smart city, Stock market, etc): e.g. computer-based medical image analysis, drug discovery, and patient diagnosis, with a focus on improving healthcare outcomes.



Are synthetic data always good?



Thank you!





Questions?



AEPIA
ASOCIACIÓN
ESPAÑOLA PARA
LA INTELIGENCIA
ARTIFICIAL

EVIA2025