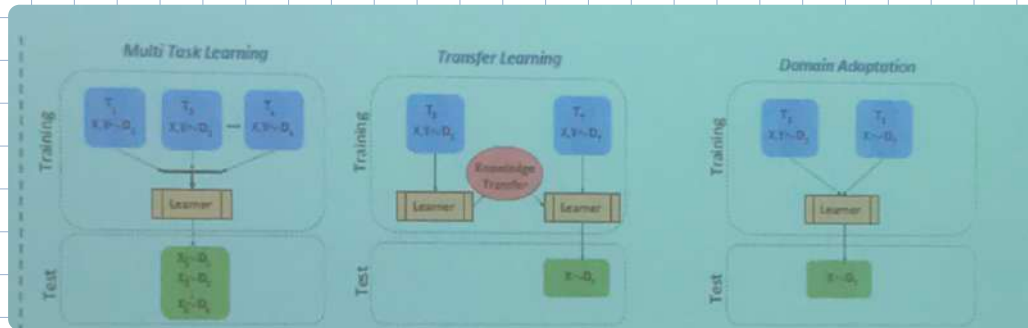


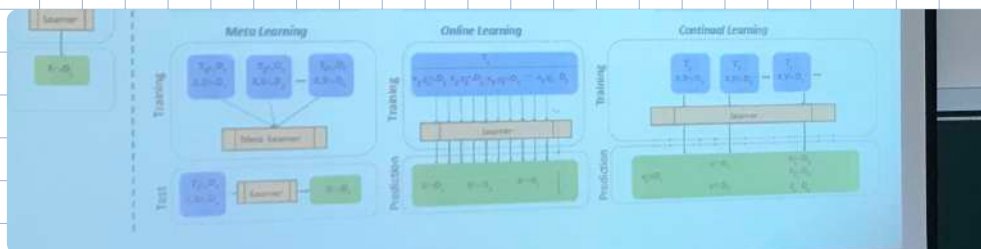
Continual Learning - Tiina Tuytelaars (KU Leuven)

$T_1 \rightarrow T_2 \rightarrow \dots \rightarrow T_N$

• Continual = Lifelong = Incremental
= Never ending



- Multi-task $\Rightarrow$ One model $\rightarrow$ multiple tasks: Each task acts as form of regularizers for the others $\rightarrow$ auxiliary tasks help a lot if drops overlap
- Transfer $\Rightarrow$ Leverage source task insights to solve the target task $\rightarrow$ From pretrained
- Domain adaptation $\Rightarrow$ Single task $\rightarrow$ shift in distr. $\Rightarrow$ sometimes no labels for target domain!
 $\hookrightarrow$ two losses: task loss + min. repr. loss between domains
 $\hookrightarrow$ adversarial loss: could be not be able to distinguish between domains
- Meta-Learning $\Rightarrow$ Task distribution Few Shot Strategy! $\rightarrow$ L-to-L
 $\hookrightarrow$ At test time apply meta-learner to specific test task $\Rightarrow$ Initialization
- Online Learning $\Rightarrow$ Streaming + no train/test time split! $\Rightarrow$
 $\hookrightarrow$ Problem: Samples are not drawn from iid fashion! $\Rightarrow$ overfits most recent data!
- Continual learning: Switch in tasks $\rightarrow$ Not all data available at train time



- ↳ learn one task after other $\Rightarrow$ no sharing of prev. data
- $\Rightarrow$ no large memory footprint $\Rightarrow$ no forgetting $\Rightarrow$ time to relax!

- ① Regularization - based $\Rightarrow$ data vs. model
- ② Rehearsal / replay - based
- ③ Network architecture - based

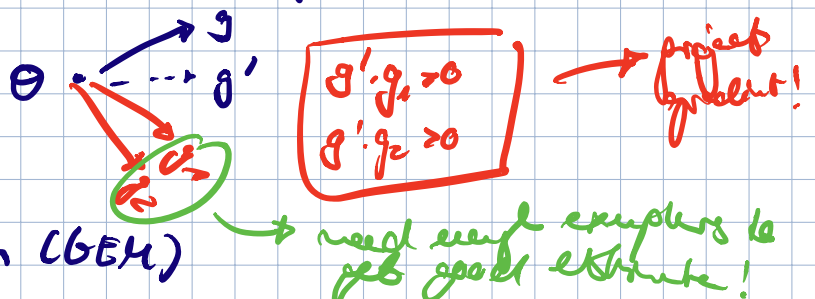
① REGULARIZATION $\Rightarrow$ Implicit Subnetworks

- Knowledge distillation loss $\Rightarrow$ preservation of responses
 - ↳ works well for related tasks Li & Hoiem (ECCV, 16)
 - ↳ Memory: store old model vs. store old predictions
- Penalize changes to important parameters $\Rightarrow$ Prior!

$$T_n: \min_{\Theta} \frac{1}{N} \sum_{n=1}^N \mathcal{L}(y_n, f(x_n, \Theta^n)) + \lambda \sum_i \mathcal{R}_i (\Theta_i^n - \Theta_i^{n-1})^2$$
 - ↳ flexibility vs. stability $\rightarrow$ how to estimate importance weights?
 - ↳ Coefficients: Elastic Weight Consolidation $\rightarrow$ class effects
 - ↳ Multiple tasks: Don't add prior for task but find a specific representative sub $\rightarrow$ memory consolidation
- Memory-aware Synapses: ensure that output does not change too much in parts of input space!
- Synaptic Intelligence $\rightarrow$ Zeh et al (ICML, 2018)

② REHEARSAL $\Rightarrow$ Representative Memory Buffer

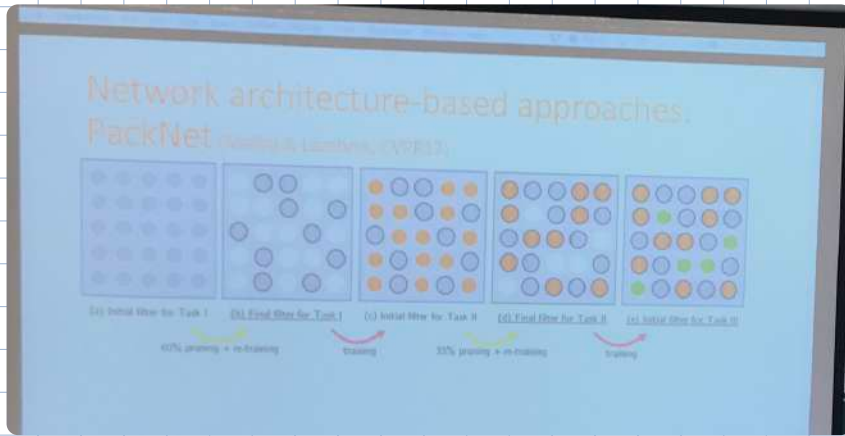
- ICARL (Rebuffi et al, 2017) $\Rightarrow$ select samples closest to mean of each class
 - ↳ Knowledge distillation loss old + new data
- Gradient Episodic Memory $\Rightarrow$ constrain gradient to improve previous tasks (GEM)
 - ↳ focus on transfer backward / forward!
- How many exemplars?
 - ↳ Fixed # per task (GEM) $\rightarrow$ need enough exemplars to get good estimate!



↳ Fixed memory (11x11) $\Rightarrow$ adapt / change later on

③ ARCHITECTURE $\Rightarrow$ Explicit Subnetworks

- PackNet (Mallya & Lazebnik, CVPR 17)



↳ Prune + retrain on sparse for task $\Rightarrow$ freeze

$\Downarrow$
Next Task 1

↳ Compression + Adaptation
↳ Awareness no forgetting!
↳ Need to know the tasks

↳ How to choose how much to prune? $\Rightarrow$ Open Question!

COMPARISON - INSIGHTS

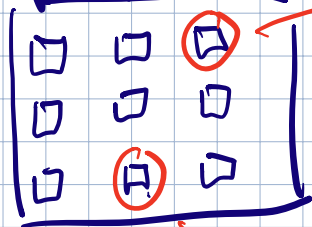
- Importance of hyperparameters: Stability vs. Flexibility
- Tiny FuzNet $\Rightarrow$ Good baseline datasets!
- PackNet implies no forgetting!
- MTS were robust than EWL
- Order of tasks does not matter too much $\rightarrow$ First tasks needs to be representative $\Rightarrow$ otherwise architecture sees lot to unlearn
- Large models $\Leftrightarrow$ More capacity $\rightarrow$ gets with the day

Looking ahead:
long term desiderata continual learning

- Constant memory
- Task agnostic
- Online learning
- Forward transfer
- Backward transfer
- Problem agnostic
- Adaptive
- No test time oracle
- Task revisiting
- Graceful forgetting

Self-Supervised Learning - Andrew Senior

- Self-Supervised = Supervision comes from the data $\rightarrow$ Proxy tasks



$\rightarrow$ CNN

$\rightarrow$ CNN



$\rightarrow$ Predict relative position to each other

$\hookrightarrow$ In order to solve you need conceptual knowledge

$\hookrightarrow$ Goal: "ImageNet" - like features without supervision

FROM IMAGES

- $\rightarrow$ PASCAL VOC $\Rightarrow$ learn embedding unsupervised and test on downstream task!
- $\rightarrow$ Problem of network learning to cheat!
 - $\hookrightarrow$ Dramatic Xliberation $\rightarrow$ Come from this! $\Rightarrow$ Use only one color channel!
- $\rightarrow$ Loss: Clustering $\rightarrow$ Give gray and let network output colored image
- $\rightarrow$ Exemplar Networks $\Rightarrow$ Classify very different types of one patch as same patch $\rightarrow$ learn invariance!
- $\rightarrow$ Multi-Task Self-Supervised Learning
 - $\hookrightarrow$ Features fed into classifiers that is then trained supervised
- $\rightarrow$ Image Transpunctics $\Rightarrow$ predict rotations
- $\rightarrow$ Jigsaw Patch Puzzle $\Rightarrow$ Borzi & Favaro, 2016
 - $\hookrightarrow$ & patch ordering / permutation
 - $\hookrightarrow$ low control complexity of predicting: How many colors to predict? // How many patches to permute?
- $\rightarrow$ Benefits from using data & sequentially more complex problem!
- $\rightarrow$ Think about in terms of Meta-learning $\Rightarrow$ Weight Init!

• FROM VIDEOS

- strong correlation in time $\Rightarrow$ Order / direction / track
- "Shuffle & Learn" $\Rightarrow$ Debias does binary decisions

T1

F1

T3

F2

T2

F3

→ Same time: have less to be info

→ Important: Need to give net a hard set of examples $\Rightarrow$ form of preprocessing to design bounds

→ main notion: Need to learn landmarks to detect action

→ Order prediction $\rightarrow$ behaviorally sort $\rightarrow$ harder than binary

→ Coldplay $\Rightarrow$ 'how of time' video $\rightarrow$ Predict flow!

→ Again be careful with decaying:

→ Zoom in in videos

→ often camera angle moves often

→ uncorrelated in black boxes

→ stabilize camera

• TEMPORAL COHERENCE OF COLOUR $\Rightarrow$ TRACUINO

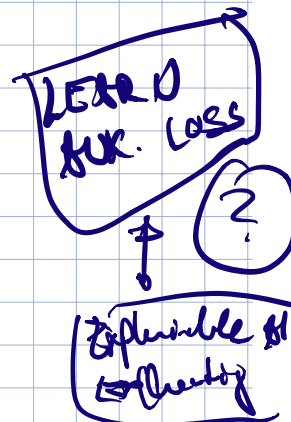
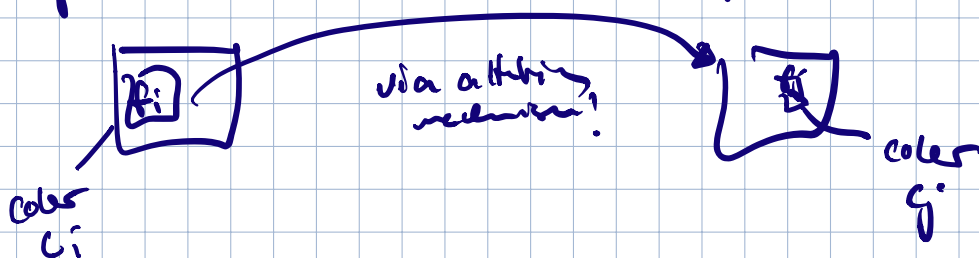
→ Self-supervised tracking $\Rightarrow$ semantic correspondence

→ Vondrick et al (ECCV, 2018)

→ Reframe Frame

$\leftrightarrow$

→ Reframe Frame



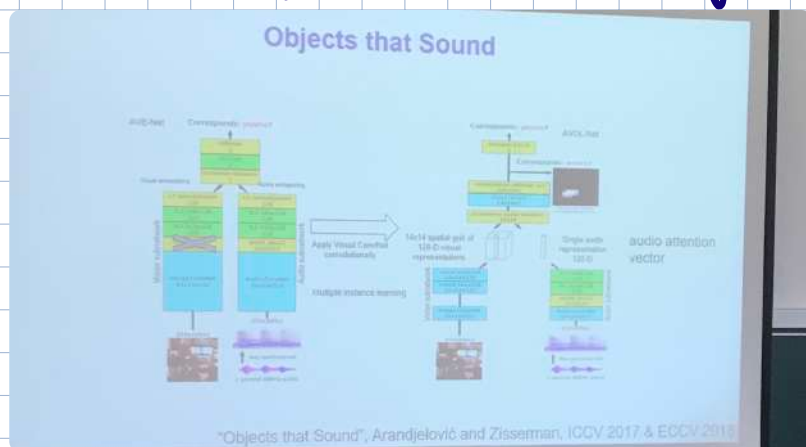
• AUDIO - VISUAL STRETCHES

① Predict synchronizability

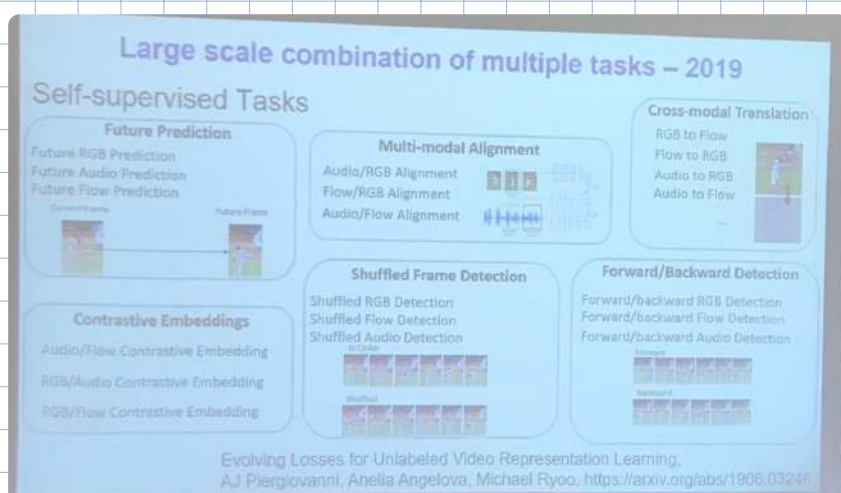
② Predict correspondence

} Proxy Tasks!

- ↳ Contrastive loss $\Rightarrow$ Min. distance positive pairs, Max. dist. neg. pairs
- ↳ Synchronization $\Rightarrow$ predict where source of audio comes from
- ↳ $\text{Image}(s) \rightarrow \text{visual NETWORK} \rightarrow \mathbb{R}^q$
- ↳ $\text{Audio} \rightarrow \text{audio NETWORK} \rightarrow \mathbb{R}^d$
- ↳ Enforce networks to learn tight embeddings!
- ↳ Correspondence $\Rightarrow$ give image + audio

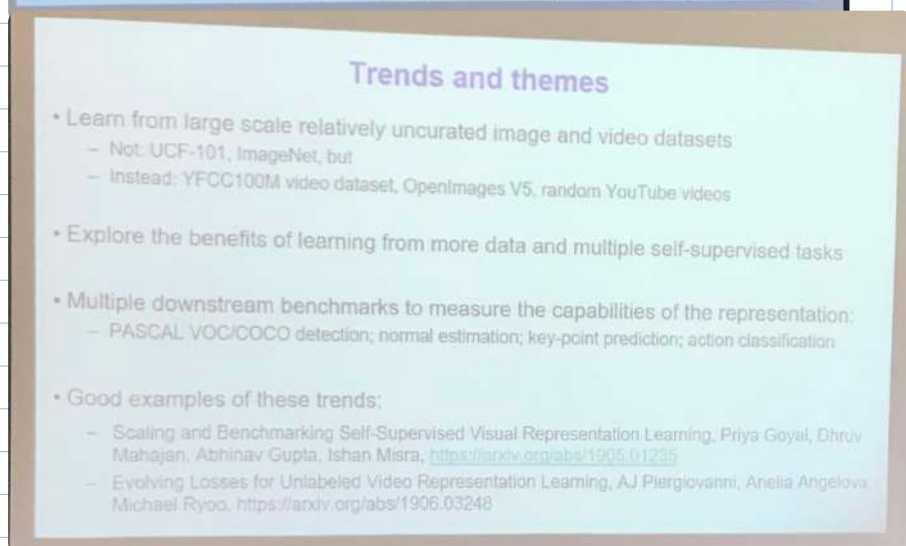


$\Rightarrow$ how to go from style frame to detection
 ↳ background vector!



$\Rightarrow$ Putting all different proxy tasks together

↳ Form of h-sket improved!
 ↳ a lot more data-points



$\Rightarrow$ Towards consistency // learning physics / geometry!

Bayesian Learning - Dmitry Vetrov $\rightarrow$ Variational Inference!

$$p(z|x) = \frac{p(x|z)p(z)}{\int p(x|z)p(z) dz} \quad \rightarrow \text{hard integration}$$

$$\begin{aligned} \log p(x) &= \int q(z) \log p(x) dz \\ &= \int q(z) \log \frac{p(x|z)}{q(z)} dz + \int q(z) \log \frac{q(z)}{p(z|x)} dz \\ &= \underbrace{\mathcal{L}(q)} + \mathcal{L}(q(z) \| p(z|x)) \\ &\quad \text{ELBO} \rightarrow \text{stochastic optimization} \end{aligned}$$

• VI: Inference = optimization! ϕ parametrizes variational distribution

$$\rightarrow p(T, w|x) = p(T|x, w) p(w) \quad \text{(DISCRIMINATIVE MODEL)}$$

$$\phi^* = \arg\max_{\phi} \int q(w|\phi) \log \frac{p(T_{tr}, w|x_{tr})}{q(w|\phi)} dw$$

$$= \underbrace{\int q(w|\phi) \log p(T_{tr}|x_{tr}, w) dw}_{\text{if we only opt. this average to converge to } \delta\text{-jet. at ML}} - \underbrace{\int q(w|\phi) \log \frac{q(w|\phi)}{p(w)} dw}_{\text{REGULARIZER!}}$$

$$= \underbrace{\left(\sum_{i=1}^n \int q(w|\phi) \log p(\tau_i|x_i, w) dw \right)}_{\text{apply minibatch!}} - \underbrace{\mathcal{L}[q(w|\phi) \| p(w)]}_{\text{still intractable } \rightarrow \text{MC estimate}} \quad \text{analytical computation!}$$

$\rightarrow$ REPARAMETRIZATION TRICK: $\int r(\epsilon) \log p(\tau_i|x_i, w(\epsilon, \phi)) d\epsilon$

$$\text{where } \epsilon \sim \mathcal{N}(0, \Sigma), w = \mu + \epsilon \circ \sigma$$

$\rightarrow$ Use MC estimation to get stochastic gradient w.r.t. $\phi = \{\mu, \sigma\}$

EEML - Day 4 Budapest

→ INVERSE RL

RL/Planning with Humans - Shua Fan (Berkeley)

- Example with over policy beta ("following optimal policy") often being pushed by human team
 - does not respect human expressed intention
 - Reward fun. has to be specified
 - Multi-agent setting $\Rightarrow$ what if other agents are humans?!

PERSON IN ENV ①

ENV - USER ②

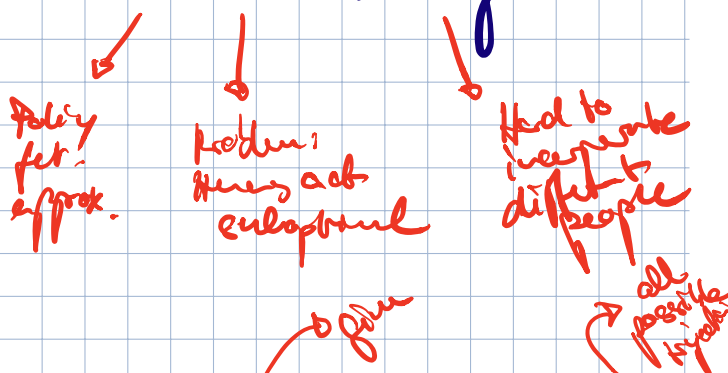
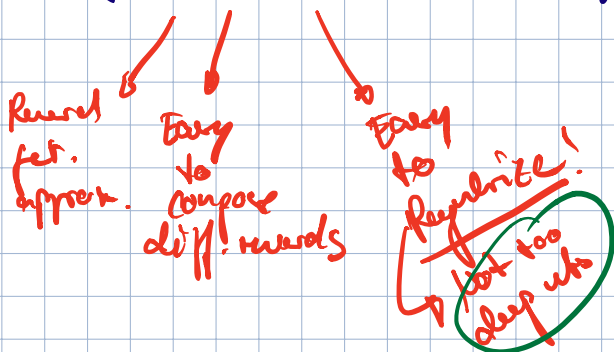
DESIGNER ③

- Value Iteration $\Rightarrow$ solve yourself if you have access to the dynamics!
- Inverse RL: From traces $\xi \rightarrow R$

→ Inverse RL

vs.

Imitation Learning



→ Problem: find $R(s,a)$ s.t. $R(\xi_0) \geq R(\xi) + \gamma$

$$R(s,a) = \theta^T \phi(s,a)$$

$$R(\xi_0) \geq R(\xi) + \gamma$$

$$\max_{\theta} [R(\xi_0) - \max_{\xi} [\theta^T \phi(\xi) + \gamma l(\xi, \xi_0)]] \quad \max_{\xi} [R(\xi) + \gamma l(\xi, \xi_0)]$$

→ optimize via GP!
 $\Rightarrow$ subproblem

max-utility planning

bonus for being far from destruction

$$\nabla_{\theta} = \phi(\xi_0) - \phi(\xi_0^*)$$

→ sum of regularizers
→ can't just do 0 rewards

- Suboptimal demonstrations $\Rightarrow$ want to be Bayesian!

$$P(\xi_0 | \theta) \propto \frac{e^{\beta \theta^T \phi(\xi_0)}}{\sum_{\xi} e^{\beta \theta^T \phi(\xi)}} \rightarrow \text{allow small prob for all trajectories}$$

$\beta = 0$: Random demonstrator $\rightarrow$ uniform distr.

$\beta < 0$: Informative demonstrator

$\beta \gg 0$: Deterministic demonstrator

- Problem: ξ are all possible trajectories $\rightarrow$ hard to compute/enumerate $\Rightarrow$ use softmax value function

- Want to do Bayesian inference: $b'(\theta) \propto b(\theta) P(\xi_0 | \theta)$
 $\hookrightarrow$ but often resort to Max LL:

$$\max_{\theta} \log \frac{e^{\beta \theta^T \phi(\xi_0)}}{\sum_{\xi} e^{\beta \theta^T \phi(\xi)}} = \beta \theta^T \phi(\xi_0) - \log \sum_{\xi} e^{\beta \theta^T \phi(\xi)}$$

$$J_{\theta} = \beta [\phi(\xi_0) - E_{\xi \sim \theta} \phi(\xi)]$$

$\rightarrow$ Exp. feature value predicted by current model

- Multi-shot problem: Robot watches former $\rightarrow$ problem: Want robot to probe other goals!

$\hookrightarrow$ Distributional shift $\Rightarrow$ need back and forth robot probes $\hookrightarrow$ learner adapts $\hookrightarrow$ robot learns