

# Pokročilé techniky hlbokého učenia

Kód kurzu: MLC\_ADV

Kurz je určený pre záujemcov o hlbšie porozumenie umelej neurónovej siete a hlavne takzvanému hlbokému učeniu. Predpokladá sa základná znalosť princípov na úrovni kurzu Úvod do strojového učenia, ktoré sa v kurze využije pre vysvetlenie pokročilejších architektur a techník. Zvláštna pozornosť bude venovaná možnostiam interpretovateľnosti modelov strojového učenia.

## Pre koho je kurz určený

Kurz je určený pre záujemcov o hlbšie porozumenie umelým neurónovým sieťam a najmä takzvanému hlbokému učeniu.

## Požadované vstupné znalosti

- Základná znalosť programovania v Pythone
- Stredoškolské znalosti lineárnej algebry, matematickej analýzy a teórie pravdepodobnosti. Predpokladá sa základné porozumenie pojmom ako vektor, matica, vektorový priestor, pravdepodobnosť, podmienená pravdepodobnosť, nezávislosť náhodných javov a znalosť násobenia matíc a derivácií funkcií.
- Znalosti strojového učenia na úrovni kurzu Úvod do strojového učenia

## Študijné materiály

Študijný materiál spoločnosti Machine Learning College.

## Osnova kurzu

- Architektúry neurónových sietí (feed-forward, rekurentné, konvolučné, generatívne, autoenkódery, Unet, GAN, attention layer)
- Optimalizátory a ich evolúcia (Steepest Gradient Descent, Stochastic Gradient Descent, Mini-Batch Gradient Descent, Nesterov Accelerated Gradient, Adagrad, AdaDelta, Adam, hľadanie učiacich koeficientov)
- Chybové funkcie a ich vlastnosti (Mean squared error, Mean absolute error, Negative Log Likelihood)
- Regularizácia neurónových sietí (Dropout, Early stopping, Data augmentation, Batch and layer normalization)
- Inicializácia neurónových sietí (Gradient vanishing problem, Zero initialization, He initialization, Xavier initialization)
- Semi-supervised learning (Pseudo Labeling, Mean-Teacher, PI-Model)
- Odhad spoľahlivosti predikcií (Logit analysis, Confidence networks)
- AutoML (automatické hľadanie hyperparametrov, grid search, Bayesian optimization, meta-learning, automatické hľadanie architektur neurónových sietí)
- Praktické príklady s knižnicou AutoKeras
- Interpretovateľnosť modelov strojového učenia (priamo interpretovateľné modely, Partial Dependence Plot, Permutation feature importance, Surrogate models, Activation Maximization, Grad-CAM)

### GOPAS Praha

Kodaňská 1441/46  
101 00 Praha 10  
Tel.: +420 234 064 900-3  
[info@gopas.cz](mailto:info@gopas.cz)

### GOPAS Brno

Nové sady 996/25  
602 00 Brno  
Tel.: +420 542 422 111  
[info@gopas.cz](mailto:info@gopas.cz)

### GOPAS Bratislava

Dr. Vladimíra Clementisa 10  
Bratislava, 821 02  
Tel.: +421 248 282 701-2  
[info@gopas.sk](mailto:info@gopas.sk)



Copyright © 2020 GOPAS, a.s.,  
All rights reserved