

Vzorové otázky na zkoušku

(<http://cmp.felk.cvut.cz/cmp/courses/recognition/Exam-questions/exam-questions.pdf>)

1. Definujte střední ztrátu strategie. Definujte význam všech použitých symbolů.

Pro strategii $q(x): X \rightarrow D$, kde X je prostor všech možných pozorování a D je prostor všech možných rozhodnutí definujeme střední ztrátu $R(q)$ jako

$$R(q) = \sum_{x \in X} \sum_{k \in K} p(x, k) W(k, q(x)),$$

kde $p(x, k)$ je sdružená pravděpodobnost toho, že budeme pozorovat x při třídě k ,

a $W(k, d): K \times D \rightarrow R$ je ztráta za rozhodnutí d při třídě k , a sčítáme přes všechny kombinace pozorování a tříd.

2. Kdy nelze rozpoznávací úlohu modelovat jako minimalizaci střední ztráty, t.j. jaká jsou omezení Bayesovské úlohy rozpoznávání?

Když neznáme pravděpodobnosti $p(x, k)$, když nelze určit smysluplná ztrátová funkce W , nebo když hraje zásadní roli čas rozhodnutí a cena pozorování (počítá se s tím, že všechna pozorování jsou předem známa).

3. Kde se stává rozhodování podle minima střední ztráty rozhodováním dle minima aposteriorní pravděpodobnosti?

Doufám, že je tam překlep a má tam být "podle maxima aposteriorní pravděpodobnosti", pak bych odpověděl, pokud máme 0/1 ztrátovou funkci, tj. za správně ohodnocenou třídu neplatíme žádnou ztrátu a za všechny druhy chyb platíme ztrátu 1.

$$\text{Parciální riziko } R(x, q) = \sum_{k \in K} p(k|x) W(k, q(x)) = \sum_{k: q(x) \neq k} p(k|x) * 1 = 1 - p(k = q(x)|x).$$

Chceme min $R(x, q)$, což dosáhneme maximalizací $p(k = q(x)|x)$, neboli za $q(x)$ volíme třídu k s největší aposteriorní pravděpodobností.

4. Co je to strategie rozhodování? Co je to Bayesovská strategie?

Strategie rozhodování je funkce $q(x): X \rightarrow D$, kde X je prostor všech možných pozorování a D je prostor všech možných rozhodnutí. Pro každé pozorování nám tedy dá určité rozhodnutí.

Bayesovská strategie je strategie q^* , která minimalizuje Bayesovské riziko $R(q) = \sum_{x \in X} \sum_{k \in K} p(x, k) W(k, q(x))$, $q^* = \operatorname{argmin}_{q: X \rightarrow D} R(q)$.

5. Odvoďte tvar diskriminační funkce pro problém klasifikace do dvou tříd A, B.

Rozdělení pravděpodobnosti ve třídách A i B jsou Gaussové se stejnou kovarianční maticí C. Střední hodnoty se liší.

Máme $p(x|A) = 1 / \sqrt{|(2\pi)^D C|} \cdot \exp(-(x - x_A)^T C^{-1} (x - x_A) / 2)$, kde $x_A \in R^D$ je střední hodnota rozdělení třídy A, x je z R^D , $C = C^T$ je z $R^{D \times D}$. Podobně pro $p(x|B)$. Hezčeji:

$$p(x|1) = \mathcal{N}(x | \mu_1, \Sigma) = (2\pi)^{-\frac{D}{2}} |\Sigma|^{-\frac{1}{2}} e^{-\frac{1}{2}(x - \mu_1)^T \Sigma^{-1} (x - \mu_1)}$$

(pro druhou třídu je rozdíl jen v μ_2)

Rozhodujeme podle poměru aposteriorních pravděpodobností:

x ohodnotíme jako A iff $p(A|x)/p(B|x) \geq 1$

x ohodnotíme jako B iff $p(A|x)/p(B|x) < 1$

Přejdeme k log odds ratio a máme:

x ohodnotíme jako A iff $\ln(p(A|x)/p(B|x)) \geq 0$

x ohodnotíme jako B iff $\ln(p(A|x)/p(B|x)) < 0$

Víme:

$$p(x|A)p(A) = p(x, A) = p(A, x) = p(A|x)p(x)$$

$p(A|x) = p(x|A)p(A)/p(x) = p(x|A)p(A)$, kde $p(x)$ zanedbáme (je něco jako konstanta, respektive dále by se ve zlomku stejně zkrátilo)

$$\ln(p(x|A)p(A)/p(x|B)p(B)) = \ln\left(\frac{1}{\sqrt{(2\pi)^B|C|}} \exp\left(-\frac{(x - x_A)^T C^{-1} (x - x_A)}{2}\right) p(A) \right) - \ln\left(\frac{1}{\sqrt{(2\pi)^B|C|}} \exp\left(-\frac{(x - x_B)^T C^{-1} (x - x_B)}{2}\right) p(B)\right) = -\frac{(x - x_A)^T C^{-1} (x - x_A)}{2} + \frac{(x - x_B)^T C^{-1} (x - x_B)}{2} + \ln(p(A)/p(B)) = -x^T C^{-1} x + 2x_A^T C^{-1} x - x_A^T C^{-1} x_A + x^T C^{-1} x - 2x_B^T C^{-1} x + x_B^T C^{-1} x_B + \ln(p(A)/p(B)) = 2(x_A - x_B)^T C^{-1} x - x_A^T C^{-1} x_A + x_B^T C^{-1} x_B + \ln(p(A)/p(B)),$$
 kde modře je lineární člen $a^T x$ a červeně je konstantní člen b . Využili jsme $C = C^T$.

Nyní máme $f(x) = a^T x + b$ a rozhodujeme:

x ohodnotíme jako A iff $f(x) \geq 0$

x ohodnotíme jako B iff $f(x) < 0$

6. Definujte rozpoznávací úlohu se třídou 'nevím' (tzv. reject option). Popište Bayesovskou strategii pro tuto úlohu.

Máme množinu tříd K . Rozhodnutí definujeme jako $D = K \cup \{\text{"nevím"}\}$.

Ztrátová funkce W je rozšířená 0/1 funkce:

$$W(k, d) = 0 \text{ iff } k = d$$

$$W(k, d) = 1 \text{ iff } k \neq d \text{ a } d \neq \text{"nevím"}$$

$$W(k, d) = \varepsilon \text{ iff } d = \text{"nevím"}$$

ε je ztráta při rozhodnutí nevím, v praktických problémech je $0 < \varepsilon < 1$.

Bayesovská strategie q^* pro tuto úlohu je:

Pro každé x nalezneme $k^* = \operatorname{argmin}_{k \in K} R(x, k) = \operatorname{argmin}_{k \in K} \sum_{i=1}^K p(i|x) = \operatorname{argmax}_{k \in K} p(k|x)$

$$q^*(x) = k^* \quad \text{iff } R(x, k^*) = 1 - p(k^*|x) < \varepsilon$$

$$q^*(x) = \text{"nevím"} \quad \text{iff } R(x, k^*) = 1 - p(k^*|x) \geq \varepsilon$$

7. Jak lze odhadnout ztrátu strategie

a) se znalostí rozdělení pravděpodobnosti ve třídách (předpokládejme jednotkovou ztrátovou matici a shodné apriorní pravděpodobnosti):

Víme rozdělení pravděpodobnosti ve třídách, neboli $p(x|k)$ pro všech K tříd k . $p(k) = 1/K$.

Máme 0/1 ztrátovou funkci (doufám, jednotková ztrátová matice vyjadřuje spíš "ziskovou" funkci, kdy $W(d, k) = 1$ při $d = k$ a $W(d, k) = 0$ při $d \neq k$).

Ztráta strategie pak je $R(q) = \sum_{x \in X} \sum_{k \in K} p(x|k)p(k)W(k, q(x))$, zjednodušeně:

$$R(q) = 1/K * \sum_{x \in X} \sum_{k: q(x) \neq k} p(x|k) = 1/K * \sum_{x \in X} (1 - p(x|q(x)))$$

b) z trénovací množiny:

Máme N vzorků (x_i, k_i) v trénovací množině.

$$R_{\text{emp}}(q) = 1/N * \sum_{i=1, \dots, N} W(k_i, q(x_i))$$

Při 0/1 ztrátové funkci to je "počet špatně klasifikovaných" / N .

8. Formulujte úlohu Neymana-Pearsona. Definujte význam všech použitých symbolů.

V úloze Neymana-Pearsona máme dvě třídy $K = \{N, D\}$, kde N značí normální stav a D značí nebezpečný stav. Množina rozhodnutí $D = K$. Určíme ϵ_D , $0 < \epsilon_D < 1$, což je maximální poměr chyb "ohodnocení D jako N ", $\epsilon_D(q) = \sum_{x: q(x) \neq D} p(x|D)$. Úloha pak vypadá:

$$q^* = \operatorname{argmin}_q \sum_{x: q(x) \neq N} p(x|N) \text{ za podmínky } \sum_{x: q(x) \neq D} p(x|D) \leq \epsilon_D,$$

kde $q: X \rightarrow D$ je rozhodovací funkce, X je množina všech možných pozorování.

Obecně hledáme práh μ pro $r(x) = p(x|N)/p(x|D)$, podle kterého pak rozhodujeme:

$$q(x) = N \text{ iff } r(x) > \mu, \quad q(x) = D \text{ iff } r(x) \leq \mu$$

U spojitých měření se místo Sum používá Integrál.

9. Formulujte minimaxní úlohu. Definujte význam všech použitých symbolů.

V úloze minmax minimalizujeme maximální chybu ϵ_k na třídách $k \in K = \{K_1, \dots, K_K\}$. Úloha pak vypadá:

$$q^* = \operatorname{argmin}_q \max_k \epsilon_k(q) = \operatorname{argmin}_q \max_k \sum_{x: q(x) \neq k} p(x|k),$$

kde $q: X \rightarrow D$ je rozhodovací funkce, X je množina všech možných pozorování.

Minmax nezáleží na apriorních pravděpodobnostech $p(k)$, odpovídá Bayesovskému klasifikátoru při neznámých $p(k)$.

U spojitých měření se místo Sum používá Integrál.

10. Formulujte Waldovu úlohu. Definujte význam všech použitých symbolů.

Waldova úloha umožňuje stanovit horní mez chyb dvou tříd tím, že zavádí možnost "nevím".

Máme $K = \{1, 2\}$, $D = \{1, 2, \text{"nevím"}\}$. Zavedeme horní meze chyb na obou třídách $0 < \epsilon < 1$.

Úloha pak zní:

$$q^* = \operatorname{argmin}_q \max_k \kappa_k \text{ za podmínky } \epsilon_1 \leq \epsilon, \epsilon_2 \leq \epsilon,$$

kde $\kappa_1 = \sum_{x: q(x) = \text{"nevím"}} p(x|1)$ a $\epsilon_1 = \sum_{x: q(x) = 2} p(x|1)$ a podobně pro κ_2 a ϵ_2 .

Obecně hledáme prahy μ_u, μ_l pro $r(x) = p(x|N)/p(x|D)$, podle kterých pak rozhodujeme:

$$q(x) = N \text{ iff } r(x) > \mu_u, q(x) = D \text{ iff } r(x) < \mu_l, q(x) = \text{"nevím"} \text{ iff } \mu_l \leq r(x) \leq \mu_u$$

U spojitých měření se místo Sum používá Integrál.

11. Jaká je optimální strategie pro úlohu Neymana-Pearsona?

Nevím, jestli chtějí zas vzoreček:

$$q^* = \operatorname{argmin}_q (\epsilon_N = \sum_{x: q(x) \neq N} p(x|N)) \text{ za podmínky } \epsilon_D = \sum_{x: q(x) \neq D} p(x|D) \leq \epsilon'_D$$

Případně slovně je obecně nejlepší vzít takovou rozhodovací funkci, která má $\epsilon_D = \epsilon'_D$,

protože s rostoucí ϵ_D klesá (či maximálně stagnuje) ϵ_N . Tato rovnost nelze vždy přímo dosáhnout u diskrétních pozorování, případně jde dosáhnout randomizací výsledků. U spojitých pozorování lze dosáhnout vždy.

12. Jaká je optimální strategie pro minimaxní úlohu?

Zase buď vzoreček:

$$q^* = \operatorname{argmin}_q \max_k (\epsilon_k(q) = \sum_{x: q(x) \neq k} p(x|k))$$

Nebo myšlenkovej pochod (<https://www.youtube.com/watch?v=r37Z209M9fM> :)):

Znovu obecně platí, že růst ϵ_k jedné třídy vede ke klesání (maximálně stagnaci) ϵ_k ostatních tříd. Dokud je ϵ_k nějaké třídy ostře vyšší než ostatní, snižujeme tuto chybu. Proto v optimálním řešení vždy platí (pro spojitý případ), že třídy s nejhorší chybou jsou dvě, $\epsilon_A = \epsilon_B$. Znovu k dosažení tohoto optima u diskrétní X je někdy třeba randomizace výsledků, u spojitě X lze vždy přímo.

13. Jaké je řešení Waldovy úlohy?

Zase buď vzoreček:

$$q^* = \operatorname{argmin}_q \max_k \kappa_k \text{ za podmínky } \epsilon_1 \leq \epsilon, \epsilon_2 \leq \epsilon$$

Nebo myšlenka, že optimum nastává, když $\epsilon_1 = \epsilon = \epsilon_2$. U Walda si nejsem tolik jistý, ale přijde mi to taky tak.

14. Ukažte, Jak je závislá chyba strategie v bayesovské úloze klasifikace do dvou tříd na změně apriorních pravděpodobností?

Pro strategii q a dvě třídy $K = \{1, 2\}$ máme rizika na třídách ϵ_1, ϵ_2 . Pokud $\epsilon_1 = \epsilon_2$, tak se $R(q)$ nemění se změnou $p(k)$. Pokud $\epsilon_1 > \epsilon_2$, tak $R(q)$ roste s rostoucí $p(1)$ a klesá s klesající $p(1)$.

15. Jaké znáte principy odhadu parametrů?

Odhad maximální věrohodnosti - Maximum likelihood estimate - MLE

Maximální aposteriorní pravděpodobnost - Maximum a posteriori probability - MAP

Bayesovské odvození - Bayesian inference - BI

16. Formulujte úlohu odhadu parametrů podle principu maximální věrohodnosti (maximum likelihood, ML).

Máme trénovací množinu $T = \{(x_i, k_i)\}_{1 \leq i \leq N}$ a známe (předpokládaný) vzorec pravděpodobnostního rozdělení s parametry Θ . Úloha zní:

$$\Theta^* = \operatorname{argmax}_{\Theta} p(T|\Theta) = \operatorname{argmax}_{\Theta} \prod_{(x, k) \in T} p((x, k)|\Theta).$$

Pro více tříd s různými rozděleními v praxi často platí, že parametry rozdělení jednotlivých tříd jsou nezávislé a proto lze řešit toto hledání maxima na každé třídě zvlášť. Zároveň často úlohu zjednoduší přechod k log-likelihood $\ln(p(T|\Theta))$.

17. Jak získáte odhad hustoty pravděpodobnosti metodou Parzenových oken?

Máme množinu T s N trénovacími body $x \in \mathbb{R}^D$. Do každého bodu x z těchto bodů vložíme jádrovou funkci $J_x(y): \mathbb{R}^D \rightarrow \mathbb{R}$. Integrál před všechny J_x musí sčítat do 1, proto pro jeden platí: $\int J_x(y) dy = 1/N$. Hustota pravděpodobnosti v bodě y pak je $\sum_{x \in T} J_x(y)$. Jako jádrové funkce se používají například vícedimenzionální normální rozdělení, nebo uniformní rozdělení na hyperkrychlích nebo hyperkoulích.

Kernel má parametr $h = \text{"šířka"}$, u normálního rozdělení je to rozptyl, u hyperkrychlí délka strany a u hyperkoulí poloměr. Musí se optimálně zvolit, např. crossvalidací pomocí log-likelihood (jako v domácím úkolu).

18. Definujte strategii rozpoznávání podle k nejbližších sousedů (k nearest neighbour rule, k-NN)

Máme trénovací množinu $T = \{(x_i, k_i)\}_{1 \leq i \leq N}$. Pro testovací bod y nalezneme k nejbližších trénovacích bodů x_i a y ohodnotíme jako té třídy, která má mezi k nejbližšími body většinu.

19. Uveďte alespoň 5 vlastností klasifikátoru dle nejbližšího souseda (1-NN)

- implementačně jednoduchý při naivním řešení se složitostí $N \cdot D$, pro zrychlení je potřeba složitější řešení přes k -D stromy nebo pravděpodobnostní nejbližší sousedy
- i při asymptotickém případě $N \rightarrow \infty$ se jeho chyba stále liší od Bayesovské, asymptotická chyba lze odhadnout $\epsilon_{\text{Bayes}} < \epsilon_{1\text{-NN}} < 2 \cdot \epsilon_{\text{Bayes}} - R/(R-1) \cdot \epsilon_{\text{Bayes}}^2$, $R = \text{pocet tríd}$
- obecně lze použít pro libovolně velké dimenze a libovolný počet trénovacích bodů, ale problémem je konstrukce vzdálenosti a případný různý význam různých dimenzí
- velká paměťová náročnost, lze snížit redukcí trénovací množiny
- sklon overfitovat - vždy klasifikujeme pouze podle třídy nejbližšího
- není třeba předpoklad o typu distribuce, ze které pochází trénovací měření

20. Jak lze urychlit klasifikátor typu 1-NN (dle nejbližšího souseda)?

Klasifikace lze urychlit použitím k -D stromů a redukcí trénovací množiny.

K-D stromy zrychlují vyhledávání podle vzdálenosti v geometrickém prostoru pomocí ukládání dat ve stromové struktuře. Při procházení stromu se pro každou větev určí minimální možná vzdálenost nejbližšího bodu ve větvi od posuzovaného bodu. Po prvním projití do listu pak můžeme ořezat větve, které jistě nemají bližší bod. Je to sublineární vyhledávání, ale složitost je funkcí dat, nejde tedy garantovat.

Dále často stačí pouhý odhad NN, protože se data většinou vyskytují ve skupinách. Redukcí trénovací množiny zas můžeme snížit N o několik řádů, třeba z tisíců na desítky. Voronoiovy diagramy - při klasifikaci stačí zjistit, v jaké buňce bod leží (to je rychlejší než počítat znovu vzdálenosti)

21. Jaký je asymptotický odhad chyby klasifikátoru 1-NN (dle nejbližšího souseda)? Co je to asymptotický odhad chyby?

Asymptotická chyba lze odhadnout $\epsilon_{\text{Bayes}} < \epsilon_{1\text{-NN}} < 2 * \epsilon_{\text{Bayes}}$, tedy neblíží se Bayesovské chybě. Asymptotický odhad chyby znamená, že určíme chybu pro $N \rightarrow \infty$.

Příklad pro dvě třídy $p(A) = p(B) = 0,5$. V určitém x máme $p(x|A) = a$, $p(x|B) = b$. Řekněme, že $a > b$. Pak $\epsilon_{\text{Bayes}} = \min(a, b) = b$. Ale $\epsilon_{1\text{-NN}} = 2 * a * b$, kde $0,5 < a \leq 1$, tedy $1 < 2 * a \leq 2$ a odtud máme $\epsilon_{\text{Bayes}} < \epsilon_{1\text{-NN}} < 2 * \epsilon_{\text{Bayes}}$.

22. Popište perceptronový algoritmus učení. Jaké má vlastnosti?

$K = \{-1, 1\}$, $X = \mathbb{R}^D$. Nejprve upravíme trénovací množinu $\{(x_i, k_i)\}_1^N$ na $\{x'_i = (1, x_i * k_i)\}_1^N$ z $\mathbb{R}^{D+1}$. Perceptron je lineární klasifikátor, nalezne nadrovinu určenou vektorem w , která oddělí data tak, že $w * x'_i > 0$, pokud jsou data lineárně separovatelná.

Algoritmus:

- 1) init $w_0 = 0$
- 2) najdi nějaké x'_i , pro které $w * x'_i \leq 0$
- 3) pokud existuje takové x'_i , polož $w_{t+1} = w_t + x'_i$ a jdi na 2)
- 4) pokud neexistuje, ukonči, $w = w_t$

Vlastnosti:

- funguje pouze pro lineárně separovatelná data, pro neseparovatelná nikdy neskončí
- nenachází ničím optimální dělicí nadrovinu, po nalezení nějaké se zastaví
- Novikoff theorem: pro lineárně separovatelná data máme jednotkový vektor u a skalár y z $\mathbb{R}^+$ tak, že $u * x'_i \geq y$, D je velikost nejdelšího vektoru z dat. Pak se Perceptron zastaví v konečném počtu kroků $t^* \leq D^2 / y^2$
- vychází z gradientní iterační metody při ztrátové funkci $J(w_t) = \sum_i -w_t * x'_i$
- pomocí dimension lifting lze oddělovat i lineárně neseparovatelná data lineární funkcí po mapování do vyšší dimenze

23. Vysvětlete, jak lze na perceptronový algoritmus učení nahlížet jako na gradientní optimalizaci.

Pokud máme ztrátovou funkci $J(w) = \sum_{x': w * x' \leq 0} -w * x'$, tak gradient $\nabla J = \sum_{x': w * x' \leq 0} -x'$.

Při gradientní optimalizaci se pohybujeme ve směru $-\nabla J$, tedy přičtení špatně klasifikovaného x' odpovídá části kroku gradientní optimalizace.

24. Pro které rozpoznávací úlohy je lineární diskriminační funkce optimálním řešením?

- lineárně separovatelná data
- normální rozdělení se stejnou kovarianční maticí a různými středními hodnotami
- nezávislé binární vstupní vektory

- multinomiální naivní Bayes (pozorování je počet výskytů jednotlivých tříd)

25. Co je to neuronová síť? Jaké znáte základní typy.

Neuronová síť je funkce $f: \mathbb{R}^D \rightarrow \mathbb{R}^K$. Je to zřetězení lineárních funkcí (maticové násobení) a nelinearit (aktivačních funkcí). Předěly mezi lineárními funkcemi se nazývají vnitřní vrstvy.

Využívá Universal approximation theorem (UAT):

- Nechť $f: \mathbb{R}^D \rightarrow \mathbb{R}$ je spojitá funkce na jednotkové hyperkrychli a $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ je spojitá, omezená, nekonstantní funkce. Pak pro libovolné $\varepsilon > 0$ existuje přirozené číslo N tak, že $|f(x) - F(x)| \leq \varepsilon$ pro všechny x z jednotkové hyperkrychle, kde $F(x) = \sum_{i=1}^{N} v_i \sigma(w_i \cdot x + b_i)$, kde w_i je z $\mathbb{R}^D$ a v_i, b_i z $\mathbb{R}$.

Díky tomu, že jednotlivé části neuronové sítě jsou diferencovatelné, lze pomocí iteračních metod získat optimální řešení (lze uvíznout v lokálním optimu, hodně záleží na inicializaci).

Typy:

- konvoluční NN: hledá lokální souvislosti, např. pro rozpoznávání obrázků
- rekurentní NN: výstup NN lze znovu použít jako vstup, např. pro rozpoznávání zvuku
- automatické enkodéry: mají úzkou vnitřní vrstvu a požadují vstup = výstup, nutí tedy NN uložit komprimovanou reprezentaci dat

26. Jaké vlastnosti má aktivační funkce v neuronové síti? Uveďte příklady.

Aktivační funkce v neuronové síti je nelineární funkce $\sigma: \mathbb{R} \rightarrow \mathbb{R}$ a v 99,9% diferencovatelná (výjimka LReLU a ReLU v 0). Zajišťuje aproximační schopnost NN (viz UAT), bez nich by NN byly pouze zřetězené Perceptrony, které by vždy šly zjednodušit do jediné vrstvy. σ je spojitá, nekonstantní a omezená funkce.

Příklady aktivačních funkcí:

- logistický sigmoid: $\sigma(x) = 1/(1 + e^{-x})$
- tanh: $\sigma(x) = (e^x - e^{-x})/(e^x + e^{-x})$
- ReLU (rectified linear unit): $\sigma(x) = \max(0, x)$
- LReLU (leaky ReLU): $\sigma(x) = \max(0, x) + \min(0, sx)$, $0 \leq s \leq 1$

27. Popište co nejpřesněji algoritmus učení dopředné neuronové sítě metodou zpětného šíření (back-propagation). Diskutujte vlastnosti.

Protože je NN pouze jedna velká diferencovatelná matematická funkce, lze použít gradientní metodu optimalizace. Pro každý parametr v NN vypočítáme derivaci výstupu NN v závislosti na něm. Pomocí chain rule jsou derivace podle parametrů dřívějších vrstev závislé na derivacích parametrů následných vrstev. Proto postupujeme odzadu. U každého parametru (nebo obecně vrstvy), vypočítáme derivaci výstupu podle parametru samotného (pro optimalizaci daného parametru) a zároveň derivaci podle vstupu vrstvy (pro propagaci do předchozích vrstev).

Vlastnosti:

- zaručuje nalezení pouze lokálního optima, velmi záleží na prvotní inicializaci parametrů
- Rychlost konvergence velmi závisí na prvotní inicializaci parametrů

28. Porovnejte vlastnosti učení přímé neuronové sítě metodou zpětného šíření (backpropagation) a učení klasifikátoru Support Vector Machine.

Učení NN i učení SVM probíhá pomocí iteračních optimalizačních metod.

Rozdíl je v tom, že funkce u SVM je konvexní (konkávní $\arg\max f(x)$ se upraví na konvexní $\arg\min -f(x)$), proto má jediné lokální minimum, neboli vždy nalezneme globálně nejlepší řešení.

Naopak funkce NN konvexní (obecně) není, proto hrozí uvíznutí v lokálním minimu nebo sedlovém bodě. O to více zde záleží na inicializaci dat, která ovlivňuje konvergenci.

29. Co je to empirická ztráta (riziko)?

Empirická ztráta je chyba na trénovací množině.

30. Co je to strukturální riziko?

Strukturální riziko je horní odhad rozdílu empirického a testovacího rizika ($R_{\text{test}} \leq R_{\text{emp}} + R_{\text{strukt}}$). Vychází ze třídy složitosti (VC dimenze) rozhodovací funkce a počtu trénovacích dat. Je více možností, jak odhadnout horní hranici testovacího rizika (různé vzorce, někdy speciálně pro určitý typ klasifikátoru) a od toho se pak odvíjí strukturální riziko - to empirické vždy známe z naší trénovací sady.

31. Co je to Vapnik-Červoněnkisova dimenze?

VC dimenze popisuje složitost rozhodovacích funkcí. Funkce s vyšší VC dimenzí mají větší tendenci overfitovat.

Význam VC dimenze je: maximální počet bodů (N), které můžeme v prostoru nějak rozmístit tak, aby každé jejich binární ohodnocení (celkově 2^N možností) bylo možné danou třídou rozhodovacích funkcí správně odlišit.

U lineární funkce je VC dimenze $D+1$, kde D je dimenze prostoru.

Pomocí VC dimenze a chyby na trénovací sadě lze s danou pravděpodobností odhadnout horní hranici chyby na testovací sadě

$$\Pr \left(R(q) \leq R_{\text{emp}}(q) + \sqrt{\frac{1}{N} \left[h \left(\log \left(\frac{2N}{h} \right) + 1 \right) - \log \left(\frac{\eta}{4} \right) \right]} \right) = 1 - \eta$$

Pro některé typy klasifikátorů (např lineární) VC dimenzi předem známe, takže se to spočítá snadno a hodí se to.

32. Jak postupujeme při návrhu klasifikátoru podle principu minimalizace strukturálního rizika?

Začínáme u tříd klasifikátorů s nižší VC dimenzí a přecházíme ke složitějším, pokud jednodušší nestačí. Tzn. u každé třídy nalezneme nejlepší rozhodovací funkci a zjistíme její klasifikační schopnosti. Nejsou-li dostačující, zkusíme složitější třídu.

33. Jaká je VC dimenze klasifikátoru 1-NN?

Nekonečno (i nekonečně mnoho bodů lze rozlišit, v testování daných bodů 1-NN nalezneme v každém bodě nejbližšího souseda z trénovacích bodů (což jsou ty samé), což je bod sám, a ohodnotí ho na jeho třídu, která je tedy vždy správná).

34. Jaká je VC dimenze orientované nadroviny?

Orientovaná nadrovina je lineární klasifikátor, takže VC dimenze je $D+1$, kde D je dimenze prostoru.

35. Co je to Support Vector Machine?

SVM je lineární klasifikátor, který oddělí lineárně separovatelná data nadrovinou, která je nejvzdálenější od nejbližších bodů. Narozdíl od Perceptronu tedy hledá v určitém smyslu neoptimálnější rozdělení.

Zároveň se dá (narozdíl od Perceptronu) zobecnit i pro lineárně neseparovatelná data přidáním ztráty za špatnou klasifikaci bodu. Další možností je dimension lifting a použití kernel tricku, díky kterému ani nemusíme znát přímo mapovací funkci.

36. Jaká minimalizační úloha se řeší při učení SVM?

V SVM získáme maximalizaci konkávní funkce:

$$\operatorname{argmax}_{\alpha} \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i^T x_j \quad \text{za podmínky } 0 \leq \alpha \leq C, \sum_i \alpha_i y_i = 0$$

Tu převedeme na minimalizaci konvexní funkce a použijeme QP optimalizaci:

$$\operatorname{argmin}_{\alpha} \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i^T x_j - \sum_i \alpha_i \quad \text{za podmínky } 0 \leq \alpha \leq C, \sum_i \alpha_i y_i = 0$$

37. Proč algoritmus učení SVM maximalizuje tzv. "margin", t.j. minimální vzdálenost bodu trénovací množiny od rozdělovací nadroviny?

Proč jako z jakého důvodu:

Protože větší margin znamená snížení strukturálního rizika a tedy lepší generalizaci na test data.

Proč jako jakým způsobem:

Protože je úloha tak formulovaná a QP optimalizací se najde globální optimum.

38. Jak řeší SVM situaci, kdy trénovací množina není lineárně separabilní?

Úloha se rozšíří o ztrátu C za špatně klasifikovaná data. Ztráta se navíc násobí vzdáleností od správného poloprostoru pro danou třídu, je to hezky vidět v primární úloze:

$$\operatorname{argmin}_{w, b, \xi} \frac{1}{2} \|w\|^2 + C \sum_{1 \leq i \leq N} \xi_i \quad \text{za podmínky } y_i(w^T x_i + b) \geq 1 - \xi_i \text{ pro } \forall i$$

Častěji se ale používá duální úloha, kde se to projeví pouze omezením alfy:

$$\operatorname{argmax}_{\alpha} \sum_{1 \leq i \leq N} \alpha_i - \frac{1}{2} \sum_{1 \leq i, j \leq N} \alpha_i \alpha_j y_i y_j x_i^T x_j \quad \text{z.p. } \sum_{1 \leq i \leq N} \alpha_i y_i = 0, 0 \leq \alpha \leq C.$$

Toto by bylo pro data, která jsou "téměř" lineárně separabilní (např. kvůli šumu) ale pár bodů spadne na špatnou stranu rozhodovací nadroviny.

Pro data "hodně" neseparabilní by se využilo SVM s kernel trikem. Výhoda je, že výpočet bude pořád cca stejně náročný, protože se tam pořád objevuje pouze výpočet skalárního součinu - není třeba přímo počítat dimension lifting. Případně jde obojí dohromady - softmargin s kernelem, pak by rozhodovací hranice mohla být např. kružnice a tolerovaly by se některé body špatně klasifikované uvnitř / vně kružnice.

39. Příznakový prostor je $\mathbb{R}^2$. Jakým způsobem byste hledali kvadratickou rozdělovací funkci v tomto prostoru? Dokážete postup zobecnit pro $\mathbb{R}^3$ či $\mathbb{R}^n$?

Kvadratická funkce má nejvýše členy řádu 2, takže bychom v $\mathbb{R}^2$ hledali koeficienty pro členy $(1, x, y, xy, x^2, y^2)$. Pro $\mathbb{R}^3$ by to bylo předpokládám $(1, x, y, z, xy, xz, yz, x^2, y^2, z^2)$. Pro obecný $\mathbb{R}^n$ pak $(1, x_i, x_i x_j)$ pro $\forall 1 \leq i, j \leq n$.

40. Příznakový prostor je $\mathbb{R}^2$. Jakým způsobem byste hledali eliptickou rozdělovací funkci v tomto prostoru? Uvažujte pouze elipsy s osami rovnoběžnými s osami souřadnicového systému.

Předpis elipsy rovnoběžné s osami je $(x - c_x)^2/a + (y - c_y)^2/b = 1$, kde $c_x, c_y \in \mathbb{R}$ a $a, b > 0$.

Máme trénovací množinu $T = \{(x_i, y_i), k_i\}_{1 \leq i \leq N}$.

Bud' bychom asi mohli mít soustavu rovnic:

$$k_i((x_i - c_x)^2/a + (y_i - c_y)^2/b) > 1 \text{ pro } \forall 1 \leq i \leq N$$

Nebo by asi šla formulovat maximalizace margin:

$$\operatorname{argmax}_{a, b, c_x, c_y} \min_i k_i((x_i - c_x)^2/a + (y_i - c_y)^2/b) \quad \text{z.p. jako v soustavě výše}$$

41. Co je to jádrová funkce? Popište učení kernel SVM. Jaké jádrové funkce se běžně používají při učení SVM.

Jádrová funkce $K(x, y)$ popisuje skalární součin mapovaných vektorů $\phi(x)^T \phi(y)$ bez znalosti mapování. Umožňuje tak použití složitých mapování (potenciálně až nekonečně-dimenzionálních - např. Gaussian kernel: $K(x_i, x_j) = \exp(-\|x_i - x_j\|^2 / 2\sigma^2)$) v dimension lifting v SVM pro lineární oddělení lineárně neseparovatelných dat.

Učení SVM probíhá stejně jako normálně, pouze funkce přijme $K(x, y)$:

$$\operatorname{argmax}_{\alpha} \sum_{1 \leq i \leq N} \alpha_i - \frac{1}{2} \sum_{1 \leq i, j \leq N} \alpha_i \alpha_j y_i y_j K(x_i, x_j) \quad \text{z.p. } \sum_{1 \leq i \leq N} \alpha_i y_i = 0, 0 \leq \alpha \leq C$$

Běžně se používají lineární: $x^T y$, polynomiální: $(1 + x^T y)^p$, exponenciální: $\exp(-\|x - y\|^2 / (2\sigma^2))$, dvouvrstvý tanh perceptron: $\tanh(ax^T y + b)$.

Změní se klasifikace:

1) Známe ϕ :

$$w = \sum_{i \in \text{sv}} \alpha_i y_i \phi(x_i), \quad b = y_{\text{sv}} - w^T x_{\text{sv}} \quad (\text{pro nějaký support vektor } x_{\text{sv}}), \quad q(x) = w^T \phi(x) + b$$

2) Neznáme ϕ :

$$b = y_{\text{sv}} - \sum_{i \in \text{sv}} \alpha_i y_i K(x_i, x_{\text{sv}}), \quad q(x) = \sum_{i \in \text{sv}} \alpha_i y_i K(x_i, x) + b$$

42. Jaké rozdílné vlastnosti mají perceptronové učení a učení SVM.

Perceptronové učení se zastaví po nalezení první nějaké oddělující nadroviny. SVM nalezne optimální nadroviny (z hlediska velikosti margin). Oba algoritmy jsou iterativní, u perceptronu máme horní mez počtu iterací (Novikoff theorem, pokud jsou data separabilní a při inicializaci na $w^0 = 0$) $t \leq D^2/\gamma^2$, kde D je $\max_i \|x_i\|$ a $\gamma: u^T x_i \geq \gamma, \|u\| = 1$.

U SVM konvergence záleží na inicializaci α v QP.

Perceptron používá přímo body x z R^D , SVM pouze jejich skalární součiny $x^T y$ z R .

43. Popište algoritmus učení metodou Adaboost. Jakou minimalizační úlohu Adaboost řeší?

AdaBoost vybírá z množiny B slabých klasifikátorů T slabých klasifikátorů $h_t(x)$, ze kterých zkonstruuje silný klasifikátor $H_T(x) = \operatorname{sign}(\sum_{1 \leq t \leq T} \alpha_t h_t(x))$.

Algoritmus:

- 1) Inicializuj $t = 1$, váhy trénovacích bodů na $D_{t=1}(i) = 1/N$
- 2) Nalezni slabý klasifikátor $h_t(x) = \operatorname{argmin}_{h \in B} e_t(h)$, kde $e_t(h) = \sum_{i: h(x_i) \neq y_i} D_t(i)$, tzn. klasifikátor s nejmenší váženou chybou
- 3) Pokud je $e_t \geq 1/2$, zastav (další klasifikátory už nezlepšují klasifikaci)
- 4) Zadej váhu nalezeného klasifikátoru $\alpha_t = 1/2 \ln((1 - e_t)/e_t)$
- 5) Updatuj váhy bodů: $D_{t+1}(i) = (D_t(i) \exp(-\alpha_t y_i h_t(x_i)))/Z_t$, kde $Z_t = 2^{-\sum_{i=1}^N (\alpha_t y_i h_t(x_i))}$ je normalizační konstanta, $D_t(i)$ je tedy rozdělení, sčítají do 1
- 6) pokud $t = T$, ukonči, jinak jdi na 2) s $t := t + 1$

AdaBoost minimalizuje horní odhad chyby $1/N \sum_i [H_T(x_i) \neq y_i] \leq \prod_{1 \leq t \leq T} Z_t$, protože:

- 1) $[H_T(x_i) \neq y_i] \iff [y_i \cdot H_T(x_i) < 0] \iff [y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x) < 0]$
- 2) $[y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x) < 0] \leq \exp(-y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x))$
- 3) $D_{T+1}(i) = D_T(i) / (N \cdot \prod_{1 \leq t \leq T} Z_t) = \exp(-y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x)) / (N \cdot \prod_{1 \leq t \leq T} Z_t)$
 $D_{T+1}(i) \cdot \prod_{1 \leq t \leq T} Z_t = \exp(-y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x)) / N$
 $\sum_i D_{T+1}(i) \cdot \prod_{1 \leq t \leq T} Z_t = \sum_i \exp(-y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x)) / N$
 $1 \cdot \prod_{1 \leq t \leq T} Z_t = \sum_i \exp(-y_i \cdot \sum_{1 \leq t \leq T} \alpha_t h_t(x)) / N$ (...protože $D_{T+1}(i)$ je rozdělení)

$$4) \quad 1/N \sum_i [H_T(x_i) \neq y_i] \leq \sum_i \exp(-y_i \sum_{1 \leq t \leq T} \alpha_t h_t(x)) / N = \prod_{1 \leq t \leq T} Z_t$$

44. Popište algoritmus k průměrů (k-means).

Máme trénovací data bez jejich tříd $T = \{x_i\}_{i=1}^N$. Chceme je rozdělit do K shluků tak, aby suma vzdáleností bodů od center shluků byla minimální:

$$(c_1, \dots, c_K)^* = \operatorname{argmin}_{(c_1, \dots, c_K)} \sum_{i=1}^N \min_k \|x_i - c_k\|^2$$

Algoritmus:

- 1) init centra c_k (náhodně nebo pomocí K++ means)
- 2) každý bod přiřaď k nejbližšímu centru: $T_k = \{x: \|x - c_k\|^2 \leq \|x - c_j\|^2 \forall j\}$
- 3) update center do těžiště shluků:

$$c_k = 1/|T_k| \cdot \sum_{x \in T_k} x \quad \text{iff } |T_k| > 0$$

$$c_k = \text{reinit} \quad \text{iff } |T_k| = 0$$
- 4) pokud se žádný shluk T_k nezměnil, ukonči, jinak jdi na 2)

Algoritmus lze zobecnit pro jakoukoliv funkci vzdálenosti $d(x, y)$, změni se pouze 2. a 3. krok. K++ means inicializace zvyšuje pravděpodobnost nalezení globálního minima. Nicméně protože jde o iterační algoritmus, stále můžeme uvíznout v lokálním extrému (lze řešit více běhy a výběrem nejlepšího řešení).

45. EM algoritmus.

EM algoritmus je iterační metoda pro odhad parametrů, když odhad nelze vyřešit analyticky (například pomocí MLE).

Máme trénovací data $T = \{x_i\}_{i=1}^N$, kde měření nemusí být kompletní. Chceme nalézt maximálně věrohodný odhad parametrů θ , které popisují pravděpodobnostní rozdělení.

Algoritmus:

- 1) init $\theta^{t=0}$ (náhodně)
- 2) vypočítat příslušnost dat x_i k jednotlivým třídám q_i :

$$q^{t+1} = \operatorname{argmax}_q L(q | \theta^t)$$

Poznámka: Příslušnost ke třídě lze vyjádřit jako pravděpodobnost. K-means dá každému x natvrdo label, kdežto EM mu přiřadí jakoby čísla $<0,1>$ tzn. s jakou pravděpodobností náleží do dané třídy. Váhy jsou tedy vlastně ty pravděpodobnosti. x , které má vyšší pravděpodobnost, že do dané třídy patří, pak více ovlivní, jak se v současné iteraci změní odhad parametrů pro tuhle třídu.
- 3) přepočíst parametry podle vah q_i :

$$\theta^{t+1} = \operatorname{argmax}_\theta L(\theta | q^{t+1})$$
- 4) jít na krok 2)

Zkouška 2020

1. (5 points) Describe the K-means algorithm (1pt). Prove that K-means converges to a local minimum (2pts). Describe the K-means++ algorithm (1pt). Give an example when employing K-means++ leads, on average, to better solution than K-means (1pt).
2. (5 points) Formulate the logistic regression problem (1 point). Derive the gradient descent learning method for this task (2,5 points).
Given the training set $T = \{(-1, -1), (-2, -1), (0, 1), (2, 1)\}$, starting point $w = (0, 0)$, and an initial step size of 1, do the first iteration of gradient descent. (1.5 points).
3. (5 points) Parameter Estimation. The reliability of light bulbs is defined by the probability density function

$$p(t) = \frac{1}{\theta \left(\frac{t}{\theta} + 1\right)^2}, \quad t \in [0, \infty],$$

where $p(t)\delta$ is the probability that a bulb fails in a short time interval $(t, t + \delta)$. That is, the probability $P(T)$ that the bulb fails in interval $(0, T)$, is

$P(T) = \int_0^T p(t)dt = 1 - \frac{1}{(\frac{T}{\theta} + 1)}$. What is the maximum likelihood estimate of the reliability parameter θ if in an experiment with two bulbs the following lifetimes have been observed:

$t_1 = 1.0, t_2 = 4.0$? (4pts)

For this setting, what is the probability that a light bulb is still OK at time 10? (1pt)

4. (5 points) Decision tree classifiers.
 - (a) (2 points) Describe the simplest top-down tree-building algorithm that selects the decision rule at a node by maximization of information gain.
 - (b) (3 points) Discuss advantages and disadvantages of decision trees.

1)

K-means: je to tu někde

Konvergence: taky tu je (v kroku 2 si nemůže žádný bod pohoršit, v kroku 3 se nemůže zhoršit globální součet, musíme zajistit konzistenci při rovnosti vzdáleností, celkově je N^k možných stavů = konečný počet)

K-means++: taky tu už je

Příklad: ten typický obdélník

2)

Viz Zkouška 2016 1)

3)

$$\underline{3)} \quad \theta^* = \arg \max_{\theta \in \mathbb{R}} p(\gamma | \theta) = \arg \max_{\theta \in \mathbb{R}} \prod_{t \in \gamma} \frac{1}{\theta (\frac{t}{\theta} + 1)^2} = \arg \max_{\theta \in \mathbb{R}} L(\theta)$$

$$l(\theta) = \ln(L(\theta)) = - \sum_{t \in \gamma} \ln \left(\theta \left(\frac{t}{\theta} + 1 \right)^2 \right) = - \sum_{t \in \gamma} \left(\ln(\theta) + 2 \cdot \ln \left(\frac{t}{\theta} + 1 \right) \right)$$

$$\frac{\partial l}{\partial \theta} = - \sum_{t \in \gamma} \left(\frac{1}{\theta} + 2 \cdot \frac{1}{\frac{t}{\theta} + 1} \cdot t \cdot \frac{-1}{\theta^2} \right) = 0 = \sum_{t \in \gamma} \left(\frac{1}{\theta} - \frac{2t}{\theta^2 (\frac{t}{\theta} + 1)} \right)$$

$$\frac{N}{\theta} = \sum_{t \in \gamma} \frac{2t}{\theta^2 (\frac{t}{\theta} + 1)} \rightarrow N \cdot \theta = \sum_{t \in \gamma} \frac{2t}{\frac{t}{\theta} + 1} = \sum_{t \in \gamma} \frac{2 \cdot t \cdot \theta}{t + \theta}$$

$\Rightarrow$ dosadíme konkrétní hodnoty $\gamma = \{1, 4\}$:

$$2 \cdot \theta = \frac{2 \cdot 1 \cdot \theta}{1 + \theta} + \frac{2 \cdot 4}{4 + \theta} \Rightarrow 2(1 + \theta)(4 + \theta) = 2(4 + \theta) + 8(1 + \theta)$$

$$\Rightarrow \theta^2 + 5\theta + 4 = 4 + \theta + 4 + 4\theta = 5\theta + 8 \Rightarrow \theta^2 - 4 = 0$$

$$\theta = \pm 2 \rightarrow \text{bereme } \underline{\underline{\theta = 2}} \text{ (aby } p(t) \text{ bylo kladné)}$$

• Pravděpodobnost, že řídlovka je OK ~ čísel 10 = $p(\text{OK}_{t=10}) = 1 - P(10)$

$$= 1 - 1 + \frac{1}{\frac{10}{2} + 1} = \frac{1}{\frac{10}{2} + 1} = \underline{\underline{\frac{1}{6}}}$$

4)

Viz Zkouška 2017 1)

Zkouška 2019

(http://cmp.felk.cvut.cz/cmp/courses/recognition/Exam-questions/test_2019-eng.pdf)

1) ?????

b) a c) nevím

1a) $p(t) = \frac{1}{\theta} e^{-\frac{t}{\theta}}$, hledáme θ_{MLE} k $\mathcal{T} = \{0, 1, 2.3, 2.7, 6.0, 12.0\}$

$$\theta_{MLE} = \operatorname{argmax}_{\theta} p(\mathcal{T}|\theta), \quad p(\mathcal{T}|\theta) = \mathcal{L}(\theta) = \prod_{t \in \mathcal{T}} \frac{1}{\theta} e^{-\frac{t}{\theta}}$$

- nájde maximum derivací loglikelihoodu $l(\theta) = \ln(\mathcal{L}(\theta)) = \sum_{t \in \mathcal{T}} \left(\ln\left(\frac{1}{\theta}\right) - \frac{t}{\theta} \right)$

$$\frac{\partial l(\theta)}{\partial \theta} = \sum_{t \in \mathcal{T}} \left(\frac{1}{\theta} \cdot \left(-\frac{1}{\theta^2}\right) - t \cdot \frac{1}{\theta^2} \cdot (-1) \right) = 0 = \sum_{t \in \mathcal{T}} \left(-\frac{\theta}{\theta^2} + \frac{t}{\theta^2} \right)$$

$$\Rightarrow \sum_{t \in \mathcal{T}} \frac{1}{\theta} = \sum_{t \in \mathcal{T}} \frac{t}{\theta^2} \quad / \cdot \theta^2$$

$$N \cdot \theta = \sum_{t \in \mathcal{T}} t \quad \Rightarrow \quad \theta_{MLE} = \frac{1}{N} \sum_{t \in \mathcal{T}} t = \text{arithmetický průměr}$$

V našem případě je $\theta_{MLE} = \frac{24}{6} = \underline{\underline{4.}}$

b)

Druhá část b) pomocí EM: ?????

1) init θ

2) $q_i = 1/\theta * e^{-t_i/\theta}$ pro $i = 1, 2, 3, 4$, $q_j = e^{-t_{\max}/\theta}$ pro $j = 5, 6$

3) $\theta = (q_1 * t_1 + q_2 * t_2 + q_3 * t_3 + q_4 * t_4 + q_5 * t_{\max} + q_6 * t_{\max}) / (\sum_{i=1..6} q_i)$

c)

$$p(\text{fail}) = \int_0^{t_{\max}} \frac{1}{\theta} \cdot e^{-t/\theta} dt = 1 - e^{-t_{\max}/\theta}$$

$$p(\text{good}) = \int_{t_{\max}}^{\infty} \frac{1}{\theta} \cdot e^{-t/\theta} dt = e^{-t_{\max}/\theta}$$

$$\theta^* = \operatorname{argmax}_{\theta} (e^{-t/\theta})^G \cdot (1 - e^{-t/\theta})^F - \text{substituujeme } \pi = e^{-t/\theta}$$

$$\pi^* = \operatorname{argmax}_{\theta} \pi^G \cdot (1 - \pi)^F = \operatorname{argmax}_{\theta} G \cdot \ln(\pi) + F \cdot \ln(1 - \pi)$$

$$\partial/\partial \pi = G/\pi - F/(1 - \pi), \text{ odtud } \pi = G/(F+G) = e^{-t/\theta}$$

$$\theta^* = t_{\max}/\ln((F+G)/G)$$

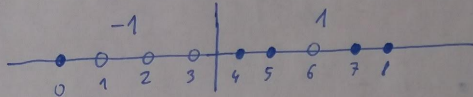
2)

2) 1) inicializujeme vektor dat na $D_i(1) = \frac{1}{N}$ pro $\forall i$ ($N=9$)

2) nalezneme nejlepší dělíč klasifikátor $h_1(x) = \operatorname{argmin}_{h \in B} \epsilon_1(h)$, kde

$$\epsilon_1(h) = \sum_{i: h(x_i) \neq k_i} D_i(1) \quad (\text{myšl. pro } t=1)$$

3) $\Rightarrow$ v našem případě bychom měli vyzkoušet $h_1(x) = \operatorname{sign}(1 \cdot x - 3.5)$ a $\epsilon_1(h) = \frac{2}{9}$



3) vypracujeme vektor klasifikátoru $h_1(x)$: $\alpha_1 = \frac{1}{2} \ln \frac{1 - \epsilon_1}{\epsilon_1} = \frac{1}{2} \ln \frac{\frac{7}{9}}{\frac{2}{9}} = 0.626$

4) upravíme vektor bodů: $D_i(2) = \frac{D_i(1) \cdot e^{-\alpha_1 \cdot h_1(x_i) \cdot k_i}}{2 \sqrt{\epsilon_1 \cdot (1 - \epsilon_1)}}$

$\Rightarrow$ pro správně klasif. body: $D_i(2) = \frac{\frac{1}{9} \cdot e^{-0.626}}{2 \sqrt{\frac{2}{9} \cdot \frac{7}{9}}} = \frac{0.535}{18 \cdot \frac{3.74}{9}} = 0.0715, i \in \{1, 2, 3, 4, 5, 6, 7\}$

$\Rightarrow$ pro špatně klasif. body: $D_i(2) = \frac{\frac{1}{9} \cdot e^{0.626}}{2 \sqrt{\frac{2}{9} \cdot \frac{7}{9}}} = \frac{1.87}{18 \cdot \frac{3.74}{9}} = 0.25, i \in \{0, 8\}$

5) pokračujeme lg. n. a hledáme $h_2(x)$ v bodu 2)

3)

3) K-NN: Máme množinu testovacích dat $\mathcal{T} = \{(x_i, k_i)\}_{i=1}^N$, kde $x_i \in \mathbb{R}^D$, $k_i \in \{1, \dots, R\}$. Pro testovací bod $y \in \mathbb{R}^D$ nalezneme K nejbližších x_i z $\mathcal{T}$ podle d a rozhodneme ho jako k_i , které má mezi K nejbližšími vítězem.

↑
funkce vzdálenosti
 $d: \mathbb{R}^D \times \mathbb{R}^D \rightarrow \mathbb{R}$

- ⊕ jednoduchá implementace (naivní)
- ⊕ funguje pro libovolný počet tříd R
- ⊕ možnosti rychlého naivní implementace: K-D stromy, výběr pravděpodobnostních NN
- ⊖ potřeba definice funkce vzdálenosti d (problém při různých "mítelkách" / různých proktu x)
- ⊖ vysoká paměťová náročnost (ke zřízení ~~redukce~~ ^{redukce} trénovací množiny)

4)

4) Maximalizujeme $IG(\mathcal{A}) = H(P) - \sum_{i \in \{0,1\}} \frac{|P_i|}{|P|} H(P_i)$, kde $\mathcal{A}$ je atribut podle kterého dělíme na množiny, tedy $\mathcal{A}$ je $i \in \{0,1\}$

$\Rightarrow$ minimalizujeme $\sum_{i \in \{0,1\}} \frac{|P_i|}{|P|} H(P_i)$, $H(X) = - \sum_{k \in \{0,1\}} \frac{|X^k|}{|X|} \ln \frac{|X^k|}{|X|}$ (X^k je množina bodů, když k v X)

- $\mathcal{A} = \text{movie length}$: $P_0 = \{4, 6\}$, $P_1 = \{1, 2, 3, 5, 7, 8\}$:

$$\frac{|P_0|}{|P|} H(P_0) + \frac{|P_1|}{|P|} H(P_1) = \frac{2}{8} \left(-\frac{1}{2} \ln \frac{1}{2} - \frac{1}{2} \ln \frac{1}{2} \right) + \frac{6}{8} \left(-\frac{3}{6} \ln \frac{3}{6} - \frac{3}{6} \ln \frac{3}{6} \right) = \frac{1}{4} \ln \frac{1}{2} - \frac{3}{4} \ln \frac{1}{2}$$

$$= -\ln \frac{1}{2} = \ln 2 \approx 0.69$$
- $\mathcal{A} = \text{country}$: $P_0 = \{2, 5, 6, 8\}$, $P_1 = \{1, 3, 4, 7\}$:

$$\frac{4}{8} \left(-\frac{2}{4} \ln \frac{2}{4} - \frac{2}{4} \ln \frac{2}{4} \right) + \frac{4}{8} \left(-\frac{2}{4} \ln \frac{2}{4} - \frac{2}{4} \ln \frac{2}{4} \right) = -\ln \frac{1}{2} = \ln 2 \approx 0.69$$
- $\mathcal{A} = \text{weather}$: $P_0 = \{1, 2, 4, 6\}$, $P_1 = \{3, 5, 7, 8\}$:

$$\frac{4}{8} \left(-\frac{3}{4} \ln \frac{3}{4} - \frac{1}{4} \ln \frac{1}{4} \right) + \frac{4}{8} \left(-\frac{1}{4} \ln \frac{1}{4} - \frac{3}{4} \ln \frac{3}{4} \right) = -\frac{1}{4} \ln \frac{1}{4} - \frac{3}{4} \ln \frac{3}{4} = \frac{1}{4} \ln 4 + \frac{3}{4} \ln \frac{4}{3}$$

$$\frac{1}{4} \ln 4 + \frac{3}{4} \ln 4 - \frac{3}{4} \ln 3 = 2 \ln 2 - \frac{3}{4} \ln 3 \approx 0.56$$

Informační zisk je maximalizován dělením podle atributu "weather" na $\{1, 2, 4, 6\}$ a $\{3, 5, 7, 8\}$.

Zkouška 2017

(http://cmp.felk.cvut.cz/cmp/courses/recognition/Exam-questions/test_2017.01.16.pdf)

1)

- 1a) 1) početní $P = 5$
- 2) maximální dílčí funkce, jejíž datům $\{P_1, \dots, P_k\}$ maximálně odpovídá informace podle IG:
- $$\max IG = H(P) - \sum_{i=1}^k \frac{|P_i|}{|P|} H(P_i), \text{ kde } H \text{ je informační entropie}$$
- 3) pro každou množinu P_i je třeba určit novou váhu a rekursivně pokračovat, dokud není 2) v každé větvi

- 1b) • Výhody: + rychlost rozhodnutí (nelineární více dimenzí vektorů)
 + nezávislé váhy, vlastnosti, méně vztahů jejich porovnání
 + méně ovládaná odlišná specifická váha
 + intuitivní a lehce interpretovatelné rozhodnutí
 + v každé větvi rozhodnutí T , kde nová váha přidá další vlastnosti pro
 lepší rozhodnutí nebo spíše kvantifikovatelnost / být
 + pro klasifikaci i regresi váhy
 + řešení souvisí na výstupní prostor rozhodnutí
- Nevýhody: - existují efektivní algoritmy, které jsou na bázi greedy algoritmu!
 $\Rightarrow$ řešení váhy optimality
 - horší overfitting (méně se používá randomizovaných lineí)

2)

- 2) SVM je lineární klasifikátor. Groupy je pro nelineární data a nachází optimální 2017
 datové matrice $\Rightarrow$ maximálně dílčí pro.
- $T = \{(x_i, y_i)\}_{i=1}^N, x_i \in \mathbb{R}^D, y_i \in \{-1, 1\}$
- primární formulace: $\arg \max_{w \in \mathbb{R}^D, b \in \mathbb{R}} \frac{1}{\|w\|}, \text{ z.p. } y_i(w^T x_i + b) \geq 1 \text{ pro } \forall i$
 $\uparrow$
 maximální vzdálenost mezi třídami, tedy body, které jsou od sebe dále
- duální formulace: $\arg \max_{\alpha} \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j x_i^T x_j, \text{ z.p. } \sum_{i=1}^N y_i \alpha_i = 0, 0 \leq \alpha_i \leq C$
- pokud chceme provést nelineární data, představíme si je jako C, která se skládá z více částí
 klasifikací a duální váha se může měnit v produkci $0 \leq \alpha_i \leq C$
- funkce v duální síle, kde pomocí QP nachází globální optimum
- Často algoritmus je obecně závislý na N. Někdy závisí na D, pokud je optimalizace rychlostí
 pomocí skalární součiny $x_i^T x_j \in \mathbb{R}$. Pokud se těmto vlastnostem nelze QP, řešení
 pak konverguje na globální inicializaci.
- Mnohdy bod z pak hodnotíme jako $d(z) = \text{sign}(w^T z + b)$, kde $w = \sum_{i \in SV} \alpha_i y_i x_i$
 je lineární kombinace support vektorů (tj. $\alpha_i \neq 0$) a $b = y_{SV} - w^T x_{SV}$ pro nějaký support.
- Těch prvků je obecně více, vol.

3) a), b)

3) a) ~~Máme danou množinu dat $\{x_i\}_{i=1}^N$, $x_i \in \mathbb{R}^D$, která chceme rozdělit do K tříd.~~
Máme danu $\{x_i\}_{i=1}^N$, $x_i \in \mathbb{R}^D$, která chceme rozdělit do K tříd.
Hledáme K center skupin tak, aby $\{c_k\}^* = \arg \min_{\{c_k\}} \sum_{i=1}^N \min_{c_k} \|c_k - x_i\|^2$
Algoritmus: 1) ~~inicializovat~~ $\{c_k\}$
2) pro každý c_k bod malých vzdáleností center
 $\Rightarrow$ vytvořit skupiny $G_k = \{x_i : \|x_i - c_k\|^2 \leq \|x_i - c_j\|^2 \text{ pro } \forall j\}$
3) průměrná ^{center} skupin do tříd: $c_k = \begin{cases} \frac{1}{|G_k|} \sum_{x \in G_k} x & \text{iff } |G_k| \neq 0 \\ \text{null} & \text{iff } |G_k| = 0 \end{cases}$
4) pokud se skupiny nemění, ukončit, jinak jdi na 2)

b) změna a provedení 2. a 3. bod: 2) $G_k = \{x_i : \|x_i - c_k\|_1 \leq \|x_i - c_j\|_1 \text{ pro } \forall j\}$
3) průměrná center skupin do mediánů

3) c), d), e)

c) zase se jím změna 2 a 3. bod. 3. bod bude stále optimalizační úloha, používá se iterativní Weiszfeldův algoritmus.

d) $I \times N \times K \times M \leftarrow$ každý prvek je jím měřící a závislosti na prvních
 $\uparrow$ $\underbrace{N \times K \times M}_{\text{je třeba iterativně upravovat}}$ každý prvek skupin s každým center
• algoritmus závislosti musí skončit v konečném čase (I je konečné číslo)

e) • inic. bod 1) se řeší: a) vyber náhodně c_1 (přibližně náhodně)
b) pro každý bod x_i vytvoř $d_i = \min_{c_k} \|c_k - x_i\|$
c) vyber c_t náhodně a x_i podle $p(x_i) = \frac{d_i^2}{\sum_{i=1}^N d_i^2}$
d) pokud $t = K$ ukončit, jinak pokračuj - buďme k p a $t++$

4) ?????

V b) zase můj odvěký nepřítel: žárovky co přežily test...

4) a) $p(t) = \theta e^{-\theta t}$, $T = \{56, 120, 424\}$, $N=3$

$$\theta_{MLE} = \operatorname{argmax}_{\theta} p(T|\theta) = \operatorname{argmax}_{\theta} \prod_{i=1}^N \theta e^{-\theta t_i} = \operatorname{argmax}_{\theta} L(\theta)$$
$$l(\theta) = \ln(L(\theta)) = \sum_{i=1}^N (\ln \theta - \theta t_i) \leftarrow \text{maximization log likelihood}$$
$$\frac{\partial l(\theta)}{\partial \theta} = \sum_{i=1}^N \left(\frac{1}{\theta} - t_i \right) = 0 \Rightarrow \frac{N}{\theta} = \sum_{i=1}^N t_i \Rightarrow \theta = \frac{N}{\sum_{i=1}^N t_i}$$

V našem případě je to $\theta = \frac{3}{56+120+424} = \frac{3}{600} = \frac{1}{200}$

Zkouška 2016

(http://cmp.felk.cvut.cz/cmp/courses/recognition/Exam-questions/test_2016.01.21-eng.pdf)

1)

Máme pravděpodobnosti $p(x|1)$, $p(x|2)$, $p(1)$ a $p(2)$. x ohodnotíme jako třídu 1 iff $p(x|1) \cdot p(1) > p(x|2) \cdot p(2)$ iff $p(x|1) \cdot p(1) / (p(x|2) \cdot p(2)) > 1$ iff $\ln(p(x|1) \cdot p(1) / (p(x|2) \cdot p(2))) > 0$.

Pokud je funkce $\ln(p(x|1) \cdot p(1) / (p(x|2) \cdot p(2)))$ lineární, problém se zjednoduší na hledání optimálních lineárně dělicích parametrů w a b . Pak:

x ohodnotíme jako 1 iff $w^T x + b > 0$ (iff $w^T x > 0$ po úpravě trénovací množiny)

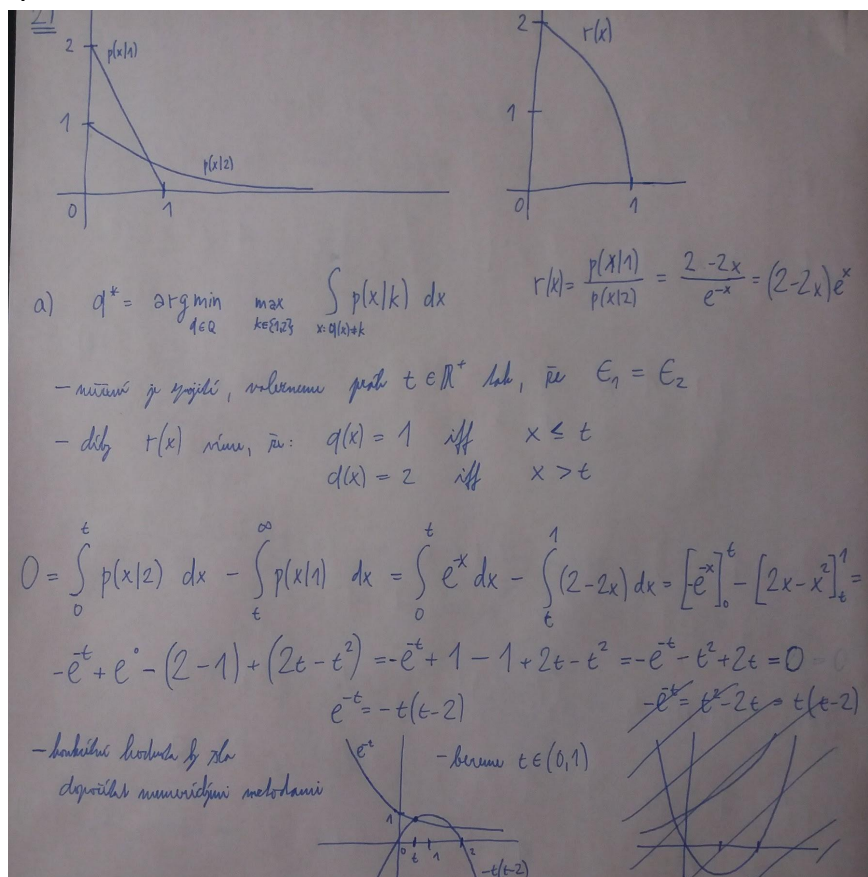
Algoritmus:

- 1) upravíme trénovací množinu $\{(x_i, y_i)\}$, $x \in \mathbb{R}^D$, $y = \pm 1$, na $\{y_i^*(x_i, 1)\}$. Iniz $w \in \mathbb{R}^{D+1}$
- 2) minimalizujeme funkci $E(w) = -l(w) = \sum_i \ln(1 + e^{-w x_i})$
Vypočítáme gradient $\nabla_w E = -\sum_i x_i^* e^{-w x_i} / (1 + e^{-w x_i})$
Udatujeme $w := w - \mu^* \nabla_w E$ (jdeme proti směru gradientu, protože minimalizujeme E)
- 3) Opakujeme krok 2)

S konkrétními hodnotami:

- 1) Získáme $T = \{(1, -1), (2, -1), (0, 1), (2, 1)\}$, $w = (0, 0)$
- 2) $\nabla_w E = -\sum_i x_i^* e^{-w x_i} / (1 + e^{-w x_i}) = -1/2((1, -1) + (2, -1) + (0, 1) + (2, 1))$ ($e^{-w x}$ je teď vždy 0)
 $\nabla_w E = -1/2(5, 0) = (-2.5, 0)$
 $w = (0, 0) - 1 \cdot (-2.5, 0) = (2.5, 0)$
- 3) Podle přednášky: Pokud $E(w_{\text{nové}}) < E(w_{\text{staré}})$, tak $\text{step} = 2 \cdot \text{step}$, jinak: $\text{step} = \text{step}/2$

2)



b) $K = \{2=D, 1=N\}$, $p(x|D) = e^{-x}$, $p(x|N) = 2-2x$ pro $x \in (0,1)$

$\Rightarrow$ novou bodinu x blíž k 0 klasifikovat jako $1=N$

$\Rightarrow \epsilon_{D=2} \leq 0.1 \Rightarrow$ máme zvýšit prahování, ke dosažení rovnosti

$$0,1 = \int_0^t e^{-x} dx = [-e^{-x}]_0^t = [-e^{-t} + e^0] = 1 - e^{-t} = 0,1$$

$$e^{-t} = 0,9$$

$$-t = \ln(0,9)$$

$$t = +0,105$$

Říšíme: $q(x) = 1 = N$ iff $x < 0,105$
 $q(x) = 2 = D$ iff $x \geq 0,105$

c) $D = \{1, 2, ?\}$, minimální velikost t_2 max pro odhad stavu 2

$\Rightarrow \epsilon_1 \leq 0.1 \Rightarrow$ bude "směrem od 1"

$$0,1 = \int_{t_2}^{\infty} p(x|1) dx = \int_{t_2}^1 (2-2x) dx + \int_1^{\infty} 0 dx = [2x - x^2]_{t_2}^1 =$$

$$0,1 = (2-1) - (2t_2 - t_2^2) = 1 - 2t_2 + t_2^2 \Rightarrow t_2^2 - 2t_2 + 0,9 = 0$$

$$t_2 = \frac{+2 \pm \sqrt{4 - 4 \cdot 1 \cdot 0,9}}{2} = 1 \pm \frac{\sqrt{0,4}}{2} = 1 \pm 0,316 \Rightarrow \text{boremu } t_2 \in (0,1)$$

Říšíme: $q(x) = 1$ iff $x < t = 0,105$
 $q(x) = 2$ iff $x > t_2 = 0,684$
 $q(x) = ?$ iff $t \leq x \leq t_2$

Uu ještě ověřit klasifikaci
 nerovnosti: K_1, K_2 :

$$K_1 = \int_{x: q(x)=?} p(x|1) dx = \int_{0,105}^{0,684} (2-2x) dx = [2x - x^2]_{0,105}^{0,684} = \underline{\underline{0,7}}$$

$$K_2 = \int_{x: q(x)=?} p(x|2) dx = \int_{0,105}^{0,684} e^{-x} dx = [-e^{-x}]_{0,105}^{0,684} = (-e^{-0,684} + e^{-0,105}) = \underline{\underline{0,316}}$$

3)

SVM:

- jen pro data z R^D
- trénovací množina může být libovolně velká, při klasifikaci se nevyužívá, pouze se z její části (tzv. support vektorů) vypočítá lineární klasifikátor $q(x) = \text{sign}(w^T x)$
- dobře generalizuje, jelikož minimalizuje VC dimenzi maximalizací margin mezi třídami dat, díky konvexnosti rozhodovací funkce a QP optimalizaci nalezne globální optimum
- při učení stačí uchovávat skalární součiny všech x_i , pro klasifikaci stačí uchovávat jediný vektor w z R^D
- učení je závislé na velikosti trénovací množiny a rychlosti konvergence QP optimalizace, která kriticky závisí na počáteční inicializaci
- vliv odchýlených hodnot lze regulovat velikostí ztráty za špatně klasifikovaný bod C , ta určuje směr "úzký margin bez špatně klasifikovaných bodů (vyšší C)" vs "širší margin se špatně klasifikovanými body (nižší C)"

DT:

- pro jakékoliv prostory, i nenumerické, s různými hodnotami/měřítky v různých dimenzích, apod.
- trénovací množina může být libovolně velká, při klasifikaci se nevyužívá (klasifikace je sublineární vůči dimenzi dat)
- žádné záruky na generalizaci/optimalitu řešení
- při učení je potřeba uchovávat veškerá data, při klasifikaci pak již DT samotný (posloupnost otázek/rozhodnutí)
- učení je závislé na složitosti stromu, respektive nic jako učení neprobíhá, strom je pouze sestaven, při stavbě stromu lze dodržovat určitá pravidla pro jeho obecně lepší klasifikaci (TDIDT)
- odchýlené hodnoty lze jednoduše odchyťovat jako speciální větve stromu

4)

Máme množinu B slabých klasifikátorů $h(x)$. Vybíráme T slabých klasifikátorů z B do silného klasifikátoru $H(x) = \text{sign}(\sum_{t=1..T} \alpha_t h_t(x))$, kde α_t jsou nezáporné váhy.

Algoritmus:

- 1) K trénovacím datům (x_i, y_i) přidáme váhu $D_1(i) = 1/N$, $x_i \in R^D$, $y_i = \pm 1$, N = počet dat, $t = 1$
- 2) Najdeme nejlepší slabý klasifikátor z B minimalizací vážené ztráty:
 $h_t = \arg\min_{h \in B} e_t(h)$, kde $e_t(h) = \sum_{i: h(i) \neq y_i} D_t(i)$
- 3) Pokud $e_t \geq 1/2$, ukonči (další klasifikátory už nepřináší zlepšení)
- 4) Váha klasifikátoru h_t je $\alpha_t = 1/2 \ln((1 - e_t)/e_t)$ ($\alpha_t > 0$ protože $e_t < 1/2$)
- 5) Přepočítáme váhy bodů: $D_{t+1}(i) = D_t(i) \cdot \exp(-\alpha_t h_t(x) y_i) / Z_t$, kde Z_t je normalizační konstanta, aby $D_t(i)$ bylo rozdělení, $Z_t = \sum_i D_t(i) \cdot \exp(-\alpha_t h_t(x) y_i) = 2 \cdot \sqrt{(e_t(1 - e_t))}$
Váhy špatně klasifikovaných se zvětší, správně klasifikovaných se sníží.
- 6) Pokud $t < T$, jdi na 2)

Odvození α_t : chceme minimalizovat Z_t :

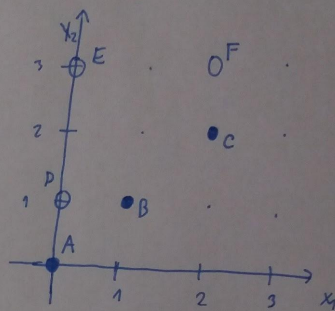
$$Z_t = \sum_i D_t(i) \cdot \exp(-\alpha_t h_t(x) y_i) = \sum_{i: h_t(x) \neq y_i} D_t(i) \cdot e^{\alpha_t} + \sum_{i: h_t(x) = y_i} D_t(i) \cdot e^{-\alpha_t}$$

$$\partial Z_t / \partial \alpha_t = \sum_{i: h_t(x) \neq y_i} D_t(i) \cdot e^{\alpha_t} - \sum_{i: h_t(x) = y_i} D_t(i) \cdot e^{-\alpha_t} = 0$$

$$\partial Z_t / \partial \alpha_t = e_t \cdot e^{\alpha_t} - (1 - e_t) \cdot e^{-\alpha_t} = 0, \text{ odtud: } e_t \cdot e^{\alpha_t} = (1 - e_t) \cdot e^{-\alpha_t}$$

$$e^{2\alpha_t} = (1 - e_t) / e_t, \text{ odtud: } \alpha_t = 1/2 \ln((1 - e_t) / e_t)$$

$$4) \mathcal{T} = \left\{ \left((0,0), 1, \frac{1}{6} \right), \left((1,1), 1, \frac{1}{6} \right), \left((2,2), 1, \frac{1}{6} \right), \left((0,1), -1, \frac{1}{6} \right), \right. \\ \left. \left((0,3), -1, \frac{1}{6} \right), \left((2,3), -1, \frac{1}{6} \right) \right\}$$



$$t=1: 1) \text{ dyby: } \epsilon_{h1m} = \frac{2}{6} \text{ pro } m=1$$

$$\epsilon_{h2m} = \frac{3}{6} \text{ pro } m=-1$$

$$\epsilon_{h3m} = \frac{3}{6} \text{ pro } m=-1$$

$$\epsilon_{h4m} = \frac{1}{6} \text{ pro } m=5 \text{ (oprotu klasifikaci D)}$$

$$2) \epsilon_{h45} < \frac{1}{2}, \quad h_1(x) = \delta(x_2 \leq 5/2)$$

$$3) \alpha_1 = \frac{1}{2} \ln\left(\frac{5/6}{1/6}\right) = \frac{1}{2} \ln(5) = 0,805$$

$$4) D_2(D) = \frac{1}{6} \cdot e^{\alpha_1} / Z_1 = \frac{2,24 \cdot 3}{6 \cdot \sqrt{5}} = \frac{1,12}{\sqrt{5}} \doteq 0,5$$

$$D_2(A, B, C, E, F) = \frac{1}{6} \cdot e^{-\alpha_1} / Z_1 = \frac{0,44 \cdot 3}{6 \cdot \sqrt{5}} = \frac{0,22}{\sqrt{5}} \doteq 0,1$$

$$Z_1 = 2\sqrt{\epsilon_1(1-\epsilon_1)} = 2\sqrt{\frac{1}{6} \cdot \frac{5}{6}} = 2\sqrt{\frac{5}{36}} = \frac{\sqrt{5}}{3}$$

$t=2: 1) \text{ dyby (distribuce je klasifikaci dle bod D):}$

$$\epsilon_{h1m} = \underline{0,2} \text{ pro } m=1$$

$$\epsilon_{h2m} = 0,3 \text{ pro } m=-1$$

$$\epsilon_{h3m} = 0,3 \text{ pro } m=7$$

$$\epsilon_{h4m} = 0,2 \text{ pro } m=1$$

$$2) \epsilon_{h11} < \frac{1}{2}, \quad h_2(x) = \delta(x_1 > 1/2) \text{ (oprotu klasifikaci A a F)}$$

$$3) \alpha_1 = \frac{1}{2} \ln\left(\frac{4/5}{1/5}\right) = \frac{1}{2} \ln(4) = 0,693$$

$$4) D_2(A, F) = 0,1 \cdot e^{\alpha_1} / Z_2 = \frac{2,5}{10,4} = \frac{1}{4}$$

$$D_2(D) = 0,5 \cdot e^{-\alpha_1} / Z_2 = \frac{0,5 \cdot 5}{2,4} = 0,3125$$

$$D_2(B, C, E) = 0,1 \cdot e^{-\alpha_1} / Z_2 = \frac{0,5 \cdot 5}{10,4} = \frac{1}{16}$$

$$Z_2 = 2\sqrt{\frac{1}{5} \cdot \frac{4}{5}} = 2\sqrt{\frac{4}{25}} = \frac{4}{5}$$

$$H_2(x) = \text{sign}\left(0,8 \cdot h_{45}(x) + 0,7 \cdot h_{11}(x)\right)$$

Zkouška 2015

(http://cmp.felk.cvut.cz/cmp/courses/recognition/Exam-questions/test_2015.01.09-eng.pdf)

1)

a)

Proběhlo celkově n hodů. Každý z nich bereme jako nezávislý, proto v každém byla pravděpodobnost q , že padne šestka, a $1-q$ že nepadne šestka. Víme, že nejdříve $n-1$ -krát nepadla, pak jednou padla, proto daný vzoreček.

b)

$$q_{MLE} = \operatorname{argmax}_q p(n|q) = \operatorname{argmax}_q (1-q)^{n-1}q = \operatorname{argmax}_q L(q)$$

$$l(q) = \ln(L(q)) = (n-1)\ln(1-q) + \ln(q)$$

$$l'(q) = -(n-1)/(1-q) + 1/q = 0$$

$$(n-1)/(1-q) = 1/q$$

$$qn-q = 1-q, \text{ odtud } qn = 1 \text{ a pak } q = 1/n$$

c)

$$q < 1/n_{\max}$$

d)

$$q_{MLE} = \operatorname{argmax}_q p(n|q) = \operatorname{argmax}_q (1-q)^{n_{\max}} = \operatorname{argmax}_q L(q)$$

$$l(q) = \ln(L(q)) = n_{\max} \ln(1-q)$$

$$l'(q) = -n_{\max}/(1-q) = 0 \dots \text{ nemá stacionární bod, musí nabývat extrému v krajním bodě:}$$

$$q \rightarrow 1 \Rightarrow l(q) = n_{\max} \ln(0) = -\infty$$

$$q \rightarrow 0 \Rightarrow l(q) = n_{\max} \ln(1) = 0$$

Což říká i intuice, že $q_{MLE} = 0$, protože čím vyšší q , tím vyšší šance, že by se objevilo, jenže ono se neobjevilo vůbec

e)

Buď to bude q_{MLE} co nám normálně vyjde, nebo krajní bod $\langle 0,1; 0,2 \rangle$, který je blíž.

$$q_{MAP} = \operatorname{argmax}_q p(q|n) = \operatorname{argmax}_q p(n|q)p(q) = \operatorname{argmax}_{q \in \{0,1\} \cup \{0,2\}} p(n|q) = \operatorname{argmax}_{q \in \{0,1\} \cup \{0,2\}} (1-q)^{n-1}q$$
$$= \operatorname{argmax}_{q \in \{0,1\} \cup \{0,2\}} L(q)$$

$$l(q) = \ln(L(q)) = (n-1)\ln(1-q) + \ln(q)$$

...

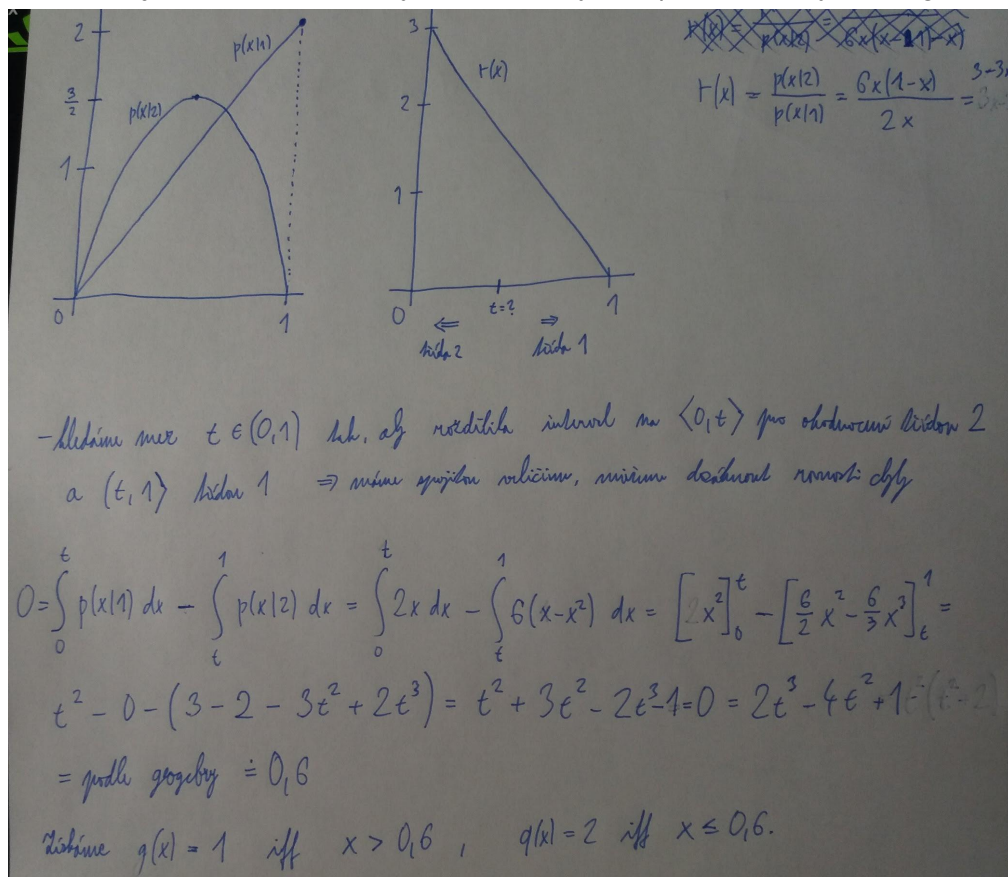
$$q_{MLE} = 1/n = 1/1 \Rightarrow q_{MAP} = 0,2$$

2)

Máme třídy $k \in \{1, \dots, K\}$. Známe $p(x|k)$ pro všechny třídy. Chceme minimalizovat největší riziko na třídě (uvažujeme 0/1 ztrátovou funkci):

$q^* = \operatorname{argmin}_q \max_k \sum_{x: q(x) \neq k} p(x|k)$ (případně integrál místo sumy u spojitých veličin)

V příkladu je předpokládám chyba a má to být: $p(x|1) = 2x$ (aby to integrovalo do 1).



3)

Máme $T = \{x_i\}_{i=1..N}$, $x_i \in \mathbb{R}^D$. Chceme je rozdělit do K shluků (tříd) tak, abychom minimalizovali sumu vzdáleností bodů od center svých tříd. Tzn. hledáme K center c_k :

$(c_1, \dots, c_K)^* = \operatorname{argmin}_{(c_1, \dots, c_K)} \sum_{i=1..N} \|x_i - c_k\|$

Algoritmus:

- 1) init $(c_1, \dots, c_K)$ (například náhodně výběrem z T nebo pomocí KM++)
- 2) každý bod zařadíme do třídy, jejíž centrum je mu nejbližší:
 $T_k = \{x_i : \|x_i - c_k\| \leq \|x_i - c_j\| \text{ pro všechna } j\}$
- 3) updatujeme centra do těžiště jejich shluků:
 $c_k = 1/|T_k| \cdot \sum_{x \in T_k} x$ (nebo reinitializujeme, když je T_k prázdné)
- 4) pokud se žádné T_k nezměnilo, ukonči, jinak jdi na 2)

a)

$C = K \cdot C_w + (\min_{((X_1, Y_1), \dots, (X_K, Y_K))} \sum_{i=1..n} \min_{k \in 1..K} \|(x_i, y_i) - (X_k, Y_k)\|) \cdot C_p$

b) + c)

Budeme iterativně postupovat pro $K = 1, 2, \dots$

Pro každé K vypočítáme C :

$$C = K \cdot C_w + \left(\min_{((X_1, Y_1), \dots, (X_K, Y_K))} \sum_{i=1..n} \min_{k \in 1..K} \|(x_i, y_i) - (X_k, Y_k)\| \right) \cdot C_p$$

$\min_{((X_1, Y_1), \dots, (X_K, Y_K))} \sum_{i=1..n} \min_{k \in 1..K} \|(x_i, y_i) - (X_k, Y_k)\|$ nalezneme upravený K-Means, který bude používat L2 euklidovskou normu bez kvadrátu:

- 1) init ($c_1, \dots, c_K$) (například náhodně výběrem z T nebo pomocí KM++)
- 2) každý bod zařadíme do třídy, jejíž centrum je mu nejbližší:
 $T_k = \{x_i : \|x_i - c_k\| \leq \|x_i - c_j\| \text{ pro všechna } j\}$
- 3) updatujeme centra pomocí Weiszfeldova algoritmu (iterativní algoritmus), tj:
 $c_k = \operatorname{argmin}_c \sum_{x \in T_k} \|x - c\|$
- 4) pokud se žádné T_k nezměnilo, ukonči, jinak jdi na 2)

Vztah K a ceny za studně je lineární a rostoucí, vztah K a ceny za trubky je nerostoucí.

Přijde mi, že optimální K je poslední K , u kterého se celková cena C nezvýší oproti předchozí (viz. Problem 10.3 b) ze sbírky)

d)

Jako v předchozím bodě, ale upravíme KM algoritmus pro L1 Manhattanskou normu:

- 1) init ($c_1, \dots, c_K$) (například náhodně výběrem z T nebo pomocí KM++)
- 2) každý bod zařadíme do třídy, jejíž centrum je mu nejbližší:
 $T_k = \{x_i : \|x_i - c_k\|_1 \leq \|x_i - c_j\|_1 \text{ pro všechna } j\}$
- 3) updatujeme centra do mediánu jejich shluku (pomocí upraveného quick sortu):
 $c_k = \operatorname{argmin}_c \sum_{x \in T_k} \|x - c\|_1$
- 4) pokud se žádné T_k nezměnilo, ukonči, jinak jdi na 2)

4)

SVM je lineární klasifikátor, nachází oddělovací nadrovinu s největším možným marginem.

$$T = \{(x_i, y_i)\}_{i=1..N}, y_i = \pm 1$$

Jeho úkol je zřejmý v původním zápisu úlohy:

$(w, b)^* = \operatorname{argmax}_{w,b} \min_i d(x_i)$ za podmínky $y_i(w^T x_i + b) > 0$, kde $d(x_i)$ je vzdálenost x_i od nadroviny

Primární úloha:

$$(w, b)^* = \operatorname{argmin}_{w,b} \frac{1}{2} \|w\|^2 \quad \text{za podmínky} \quad y_i(w^T x_i + b) \geq 1$$

Používá se ale duální forma úlohy:

$$\alpha^* = \operatorname{argmax}_\alpha \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i^T x_j \quad \text{za podmínky} \quad \alpha \geq 0, \sum_i y_i \alpha_i = 0$$

Nenulová α_i pak určují support vectory SV, které ovlivňují řešení:

$$w = \sum_{i \in \text{SV}} \alpha_i y_i x_i, \quad b = y_{\text{sv}} - w^T x_{\text{sv}}$$

Rozhodujeme pak $q(z) = \operatorname{sign}(w^T z + b)$.

Funkce v argmax je konvexní, maximalizaci tedy řešíme QP, které nám dá globální maximum. Rychlost učení velmi závisí na inicializaci.

SVM dobře generalizují, používají pouze skalární součiny dat, tzn. že jsou nezávislé na vstupní dimenzi, umí pracovat i s neseparabilními daty, při klasifikaci stačí znát pouze w a b , lze zobecnit pro oddělení nelineárně separovatelných dat pomocí dimension lifting a kernel triku.

SVM se složitě zobecňují pro oddělení více tříd a nehodí se do inkrementálního učení.

Neseparovatelná data:

Upravíme úlohy přidáním ztráty za špatnou klasifikaci bodu C:

Primární úloha:

$$(w, b)^* = \operatorname{argmin}_{w,b} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad \text{za podmínky} \quad y_i(w^T x_i + b) \geq 1 - \xi_i, \xi_i \geq 0$$

Duální úloha:

$$\alpha^* = \operatorname{argmax}_\alpha \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i^T x_j \quad \text{za podmínky} \quad 0 \leq \alpha \leq C, \sum_i y_i \alpha_i = 0$$

Při použití dimension lifting $\phi(x): \mathbb{R}^D \rightarrow \mathbb{R}^M$ nemusíme používat/znát přímo předpis mapování. Stačí znát kernel funkci $K(x, y) = \phi(x)^T \phi(y)$. Tou lze nahradit například i mapování do teoreticky nekonečně dimenzionálních prostorů apod.

Učení SVM probíhá stejně, pouze duální úloha přijme kernel funkci:

$$\alpha^* = \operatorname{argmax}_\alpha \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j K(x_i, x_j) \quad \text{za podmínky} \quad 0 \leq \alpha \leq C, \sum_i y_i \alpha_i = 0$$

Klasifikace má dva způsoby:

1) použijeme mapování $\phi(x)$:

$$\text{Získáme } w = \sum_{i \in \text{SV}} \alpha_i y_i \phi(x_i), \quad b = y_{\text{sv}} - w^T \phi(x_{\text{sv}}) \text{ a klasifikujeme}$$

$$q(z) = \operatorname{sign}(w^T \phi(z) + b)$$

To znamená, že musíme znát ϕ a klasifikace je zpomalena oproti lineární variantě výpočtem $\phi(z)$

2) použijeme kernel $K(x, y)$:

$$b = y_{\text{sv}} - \sum_{i \in \text{SV}} \alpha_i y_i K(x_i, x_{\text{sv}})$$

$$q(z) = \operatorname{sign}(b + \sum_{i \in \text{SV}} \alpha_i y_i K(x_i, z))$$

To znamená, že si musíme pro klasifikaci pamatovat α_i, y_i, x_i pro support vektory a klasifikace je zpomalena výpočtem $\sum_{i \in \text{SV}} \alpha_i y_i K(x_i, z)$

Sbírka příkladů

(https://cw.fel.cvut.cz/b201/media/courses/be5b33rpz/labs/rpz_exercise_book.pdf)

1.1

Protože $p(A|B) \cdot p(B) = p(A \cap B) = p(B \cap A) = p(B|A) \cdot p(A)$

1.2

a)

$$p(\text{rain}) = 0,27$$

$$p(\text{no rain}) = 0,73$$

$$p(1) = 0,4$$

$$p(2) = 0,4$$

$$p(3) = 0,15$$

$$p(5) = 0,05$$

b)

$$p(1 \text{ nebo } 2 | \text{rain}) = (p(1 | \text{rain}) + p(2 | \text{rain})) / p(\text{rain}) = 0,14 / 0,27$$

1.3

Cancer = {yes, no}, Test = {positive, negative}. Víme:

$$p(+, y) = 0,98, p(-, y) = 0,02, p(-, n) = 0,97, p(+, n) = 0,03, p(y) = 0,001$$

a)

$$p(y, +) = p(+, y) \cdot p(y) / (p(+, y) \cdot p(y) + p(+, n) \cdot p(n)) = 0,98 \cdot 0,001 / (0,98 \cdot 0,001 + 0,03 \cdot 0,999) = 0,0316$$

b)

$$p(n, -) = p(-, n) \cdot p(n) / (p(-, n) \cdot p(n) + p(-, y) \cdot p(y)) = 0,97 \cdot 0,999 / (0,97 \cdot 0,999 + 0,02 \cdot 0,001) = 0,99998$$

1.4

$D = (y, n)$, pro každé x můžeme spočítat parciální riziko $R(x, q) = \sum_k p(k|x) \cdot W(k, q(x))$

a)

$$R(x, q_{\text{yes}}) = 0,4 \cdot 0 + 0,2 \cdot C + 0,2 \cdot C + 0,2 \cdot C = 0,6 \cdot C$$

$$R(x, q_{\text{no}}) = 0,4 \cdot C + 0,2 \cdot 0 + 0,2 \cdot 0 + 0,2 \cdot 0 = 0,4 \cdot C$$

Nižší Bayesovské riziko je pro volbu "no"

b)

$$R(x, q_{\text{yes}}) = 0,4 \cdot 0 + 0,2 \cdot C/2 + 0,2 \cdot C/2 + 0,2 \cdot C/2 = 0,6 \cdot C/2 = 0,3 \cdot C$$

$$R(x, q_{\text{no}}) = 0,4 \cdot C + 0,2 \cdot 0 + 0,2 \cdot 0 + 0,2 \cdot 0 = 0,4 \cdot C$$

Nižší Bayesovské riziko je pro volbu "yes"

1.5

a)

Počítáme parciální riziko $q^*(x) = \operatorname{argmin}_q R(q, x)$

$$R(q, x) = \sum_{k \in K} p(k|x) \cdot W(k, q(x))$$

$$= R(d, x) = \sum_{k \in K} p(k|x) \cdot W(k, d) = \sum_{k \neq d} p(k|x) \cdot 1 = 1 - p(d, x)$$

$$\operatorname{argmin}_q R(q, x) = \operatorname{argmin}_d (1 - p(d, x)) = \operatorname{argmax}_d p(d, x)$$

b)

$D = K = (0,1)$. Bayesovské rozhodování $q(x)$ lze vyvodit ze zlomku aposteriorních pravděpodobností.

$$q_0(x) = 0, \text{ parciální riziko } R(x, q_0) = \sum_{k \in d} p(k|x) \cdot 1 = p(1|x)$$

$$q_1(x) = 1, \text{ parciální riziko } R(x, q_1) = \sum_{k \in d} p(k|x) \cdot 1 = p(0|x)$$

Bereme menší riziko, což lze vyčíst už předem ze zlomku $p(0|x)/p(1|x)$.

Pokud $p(0|x)/p(1|x) \geq 1$, volíme 0, jinak volíme 1.

Po dosazení $p(0|x) = p(x|0)p(0)/p(x)$ a $p(1|x) = p(x|1)p(1)/p(x)$ získáme:

$$p(x|0)p(0) / p(x|1)p(1) \geq 1 \text{ a odtud } p(x|0)/p(x|1) \geq p(1)/p(0) = \text{const} = \Theta$$

c)

$$p(x|0) = 1/(\sqrt{2\pi}) \cdot \sigma_0 \cdot \exp(-(x - \mu_0)^2/(2\sigma_0^2)), \text{ obdobně pro } p(x|1).$$

$$p(0|x)/p(1|x) = p(x|0) \cdot p(0) / (p(x|1) \cdot p(1)) =$$

$$= \sigma_1 \cdot p(0) / (\sigma_0 \cdot p(1)) \cdot \exp(-(x - \mu_0)^2/(2\sigma_0^2) + (x - \mu_1)^2/(2\sigma_1^2))$$

$$p(0|x)/p(1|x) < 1 \text{ je ekvivalentní s } \ln(p(0|x)/p(1|x)) < 0,$$

$$\text{neboli } q(x) = 0 \text{ iff } p(0|x)/p(1|x) \geq 1 \text{ iff } \ln(p(0|x)/p(1|x)) \geq 0$$

$$\ln(p(0|x)/p(1|x)) = 0 = -(x - \mu_0)^2/(2\sigma_0^2) + (x - \mu_1)^2/(2\sigma_1^2) + \ln(\sigma_1 \cdot p(0) / (\sigma_0 \cdot p(1)))$$

$$0 = -\sigma_1^2 \cdot (x - \mu_0)^2 + \sigma_0^2 \cdot (x - \mu_1)^2 + 2 \cdot \sigma_0^2 \cdot \sigma_1^2 \cdot \ln(\sigma_1 \cdot p(0) / (\sigma_0 \cdot p(1)))$$

$$0 = -\sigma_1^2 \cdot (x^2 - 2\mu_0 \cdot x + \mu_0^2) + \sigma_0^2 \cdot (x^2 - 2\mu_1 \cdot x + \mu_1^2) + 2 \cdot \sigma_0^2 \cdot \sigma_1^2 \cdot \ln(\sigma_1 \cdot p(0) / (\sigma_0 \cdot p(1)))$$

$$0 = (\sigma_0^2 - \sigma_1^2) \cdot x^2 + 2 \cdot (\sigma_0^2 \cdot \mu_1 - \sigma_1^2 \cdot \mu_0) \cdot x + \sigma_0^2 \cdot \mu_1^2 - \sigma_1^2 \cdot \mu_0^2 + 2 \cdot \sigma_0^2 \cdot \sigma_1^2 \cdot \ln(\sigma_1 \cdot p(0) / (\sigma_0 \cdot p(1))) = ax^2 + bx + c$$

1.6

$$\text{Oblacnost } O = \{1,2,3,4\}$$

Vypočítal jsem:

$$p(\text{rain}|1) = 0,05, p(\text{rain}|2) = 0,3, p(\text{rain}|3) = 0,6, p(\text{rain}|4) = 0,8$$

Pro rozhodnutí mi vysla rizika (p = pravdepodobnost deště)

$$R(\text{umbrella}) = 0 + 5(1 - p) = 5 - 5p$$

$$R(\text{no umbrella}) = 10p - 2(1 - p) = 12p - 2$$

$$R(100) = 5p - 0 = 5p$$

Vysla mi optimalni reseni:

$$O = 1: \text{no umbrella (s ocekavanou ztratu -1,4)}$$

$$O = 2: 100 \text{ (s ocekavanou ztratu 1,5)}$$

$$O = 3: \text{umbrella (s ocekavanou ztratu 2)}$$

$$O = 4: \text{umbrella (s ocekavanou ztratu 1)}$$

1.7

a)

3 bity musely být 000, 011, 101 nebo 110. Pro dane pravdepodobnosti mame:

$$000 = 0,7 \cdot 0,6 \cdot 0,3 = 0,126$$

$$011 = 0,7 \cdot 0,4 \cdot 0,7 = 0,196 \text{ s pravdepodobností } 0,196/0,484 = 0,405$$

$$101 = 0,3 \cdot 0,6 \cdot 0,7 = 0,126$$

$$110 = 0,3 \cdot 0,4 \cdot 0,3 = 0,036$$

b)

$$R(\text{request again}) = 0,405 \cdot \text{“zbytečný request”} = 0,405 \cdot C$$

$$R(\text{skip}) = 0,595 \cdot \text{“přehlédnout chybu”} = 0,595 \cdot 100 \cdot C$$

Určitě se vyplatí request again

1.8

Máme $p(k) = 1/3$, pravděpodobnosti $p(k|x) = p(x|k) \cdot p(k) = p(x|k)/3$:

$$p(1|0) = 0,399/3 = 0,133, \quad p(2|0) = 0,199/3 = 0,066, \quad p(3|0) = 0,065/3 = 0,022$$

$$p(1|1) = 0,242/3 = 0,081, \quad p(2|1) = 0,176/3 = 0,059, \quad p(3|1) = 0,121/3 = 0,040$$

(přepsal jsem se a mám třídy 0,1,2 místo 1,2,3)

a) $W = 0/1, x = 0$

$$R(x = 0, q(x) = 0) = 0,133 \cdot 0 + 0,066 \cdot 1 + 0,022 \cdot 1 = 0,088$$

$$R(x = 0, q(x) = 1) = 0,133 \cdot 1 + 0,066 \cdot 0 + 0,022 \cdot 1 = 0,155$$

$$R(x = 0, q(x) = 2) = 0,133 \cdot 1 + 0,066 \cdot 1 + 0,022 \cdot 0 = 0,199$$

Bereme nejmenší parciální riziko, takže $q(x) = 0$

b) $W = 0/1, x = 1$

$$R(x = 1, q(x) = 0) = 0,081 \cdot 0 + 0,059 \cdot 1 + 0,040 \cdot 1 = 0,099$$

$$R(x = 1, q(x) = 1) = 0,081 \cdot 1 + 0,059 \cdot 0 + 0,040 \cdot 1 = 0,121$$

$$R(x = 1, q(x) = 2) = 0,081 \cdot 1 + 0,059 \cdot 1 + 0,040 \cdot 0 = 0,140$$

Bereme nejmenší parciální riziko, takže $q(x) = 0$

c) $W = 0/1/2, x = 0$

$$R(x = 0, q(x) = 0) = 0,133 \cdot 0 + 0,066 \cdot 2 + 0,022 \cdot 1 = 0,154$$

$$R(x = 0, q(x) = 1) = 0,133 \cdot 2 + 0,066 \cdot 0 + 0,022 \cdot 1 = 0,288$$

$$R(x = 0, q(x) = 2) = 0,133 \cdot 1 + 0,066 \cdot 1 + 0,022 \cdot 0 = 0,199$$

Bereme nejmenší parciální riziko, takže $q(x) = 0$

b) $W = 0/1/2, x = 1$

$$R(x = 1, q(x) = 0) = 0,081 \cdot 0 + 0,059 \cdot 2 + 0,040 \cdot 1 = 0,158$$

$$R(x = 1, q(x) = 1) = 0,081 \cdot 2 + 0,059 \cdot 0 + 0,040 \cdot 1 = 0,202$$

$$R(x = 1, q(x) = 2) = 0,081 \cdot 1 + 0,059 \cdot 1 + 0,040 \cdot 0 = 0,140$$

Bereme nejmenší parciální riziko, takže $q(x) = 2$

Pravděpodobnost špatného rozhodnutí při $W = 0/1, x = 1$:

$$\text{Rozhodli jsme pro } d = 1. \quad p(k \neq d) = (p(2|x) + p(3|x)) / (p(1|x) + p(2|x) + p(3|x)) = 0,4375.$$

1.9

a)

$$K = \{0, 1, \dots, K\}, \quad X = \{1, 2, \dots, K\}, \quad D = \{\text{yes, no}\}$$

b)

Nejsem si jistý, ale tak snad, dává mi to smysl.

$$p(k \geq K/2) = q^{K/2}$$

$$p(k \geq K/2 | x) = q^{K/2 - x} \quad \text{iff } x < K/2$$
$$= 1 \quad \text{iff } x \geq K/2$$

c)

$$q(x) = n \quad \text{iff } x < K/2 \text{ \&\& } q^{K/2 - x} < 1/2$$
$$= y \quad \text{iff } x < K/2 \text{ \&\& } q^{K/2 - x} \geq 1/2$$
$$= y \quad \text{iff } x \geq K/2$$

2.1

Uděláme si poměr $p(x|M)/p(x|F)$:

0,056 (11)	0,034 (12)	0,117 (9)	0,647 (7)
0,065 (10)	0,237 (8)	2,814 (3)	2,235 (4)
2 (5)	1,75 (6)	inf (1)	inf (2)

Nejdříve vycházíme z toho, že vše klasifikujeme jako ženy. Postupně klasifikujeme jako muže x od nejvyšších poměrů, při tom kontrolujeme abychom nepřesáhli 0,2 u součtu $p(x|D)$ těchto x .

Získáme:

$$\epsilon_D = 0 + 0 + 0,145 + 0,017 + 0,001 + 0,008 + 0,017 = 0,188$$

$$\epsilon_N = 0,011 + 0,005 + 0,011 + 0,005 + 0,071 = 0,103$$

Pak lze určit mez z (0,237; 0,647), např. $\Theta = 0,4$ a rozhodovat:

x ohodnotit jako muže iff $p(x|M)/p(x|F) > \Theta$

x ohodnotit jako ženu iff $p(x|M)/p(x|F) < \Theta$

Buď skončíme takto, nebo lze optimum ještě vylepšit randomizací. Máme možnost ještě ohodnotit $0,2 - 0,188 = 0,012$ žen špatně $\Rightarrow$ další x na řadě (8) má u žen $p(x|D) = 0,299$.

Proto u x (8) můžeme ohodnotit $0,012/0,299 = 0,04 = 4\%$ vzorků jako muže, čímž získáme:

$$\epsilon_D = 0 + 0 + 0,145 + 0,017 + 0,001 + 0,008 + 0,017 + 0,04 \cdot 0,299 = 0,2$$

$$\epsilon_N = 0,011 + 0,005 + 0,011 + 0,005 + 0,96 \cdot 0,071 = 0,100$$

2.2

Nejhorší případ úspěšnost strategie q je, když třída k , která má vyšší $\epsilon_k(q)$ (tzn. při trénování bylo $p(k) < p(\text{druhá } k)$), bude mít při testování $p(k) = 1$.

a)

$$R(q) = \sum_{k \in K} \int_{x \in X} p(x,k) \cdot W(k, q(x)) = \sum_{k \in K} p(k) \int_{x \in X} p(x|k) \cdot W(k, q(x))$$

$$R(q, \pi) = \pi \cdot \int_{x \in X} p(x|1) \cdot W(1, q(x)) + (1 - \pi) \cdot \int_{x \in X} p(x|2) \cdot W(2, q(x)) = \pi \cdot \epsilon_1 + (1 - \pi) \cdot \epsilon_2$$

$$R(q, \pi) = \pi \cdot (\epsilon_1 - \epsilon_2) + \epsilon_2, \text{ kde } \epsilon_k \text{ je ztráta na třídě } k$$

$$R(q) \text{ roste s } \pi \text{ iff } \epsilon_1 > \epsilon_2, R(q) \text{ je konst. iff } \epsilon_1 = \epsilon_2, R(q) \text{ klesá s } \pi \text{ iff } \epsilon_1 < \epsilon_2$$

b)

$R(q, \pi)$ je lineární funkce π , extrému nabývá v krajním bodě definičního oboru $<0, 1>$

(v případě konstantní funkce je nezávislá na změny π a extrém je kdekoliv v $<0, 1>$)

c)

$$\max_{\pi \in \{0,1\}} (\epsilon_2 + (\epsilon_1 - \epsilon_2) \cdot \pi) = \max\{\epsilon_1, \epsilon_2\}$$

Po úpravě na 0/1 loss máme $\epsilon_1 = \int_{x: g(x) \neq 1} p(x|1)$. Chceme minimální bayesovské riziko, neboli minimalizujeme max z těchto integrálů, což je minimax.

d)

$R(q)$ je konstantní když $\epsilon_1 = \epsilon_2$. To je přesně když se chyba na obou třídách rovná, neboli jsme dosáhli minimálního maxima (další snížení chyby na jedné ze tříd by znamenalo zvýšení druhé).

2.3

a)

$q^* = \operatorname{argmin}_q \epsilon_q(N)$ za podmínky $\epsilon_q(D) \leq \epsilon$,

kde $\epsilon_q(N) = \sum_{x: q(x) = D} p(N|x)$, $\epsilon_q(D) = \sum_{x: q(x) = N} p(D|x)$ a ϵ je určená mez.

b)

Hledáme mez θ poměru $p(x|N)/p(x|D) = r(x)$ tak, abychom rozhodovali

x ohodnotit jako N iff $r(x) \geq \theta$

x ohodnotit jako D iff $r(x) < \theta$

$r(x) = 1/(x + 0,5)$ je klesající funkce na $<0, 1>$, pozorování blíže k 0 budeme označovat za N , blíže k 1 za D . Chceme abychom ohodnotili špatně 10% D stavů, neboli

$$0,1 = \int_{0 < t < H} (t + 0,5) dt = [0,5t^2 + 0,5t]_0^H = 0,5H^2 + 0,5H - 0$$

$$0 = H^2 + H - 0,2, H_{1,2} = -0,5 \pm 0,67, \text{ ale chceme } H \text{ z } <0, 1>, \text{ takže } H = 0,17$$

θ pak získáme jako $1/(0,17 + 0,5) = 1,5$, takže hodnotíme:

x ohodnotit jako N iff $r(x) \geq \theta = 1,5$

x ohodnotit jako D iff $r(x) < \theta = 1,5$

2.4

a)

$q^* = \operatorname{argmin}_q \max_k \epsilon_q(k)$, kde $\epsilon_q(k) = \int_{x: q(x) \neq k} p(x|k)$ je chyba na k -te tride

b)

Máme jednu mez θ pro $r(x) = p(x|1)/p(x|2)$. Po proložení grafu $r(x)$ konstantním $x = \theta$ máme ale dva průsečíky, které dělí x na 3 intervaly $<-1, t_1>$, $<t_1, t_2>$, $(t_2, 1>$, navíc díky symetrii víme, že $t_1 = -t_2$.

c)

U minimaxu u dvou tříd se spojitými pozorování víme, že optimum je v situaci kdy $\epsilon_q(1) = \epsilon_q(2)$. Díky symetrii se navíc může omezit na interval $<0, 1>$, kde nalezneme mez t , kterou poté promítneme jako $-t$ do $<-1, 0>$.

Na $<0, 1>$ máme $p(x|1) = x$, $p(x|2) = 1 - x$. $r(x) = p(x|1)/p(x|2) = x/(1 - x)$ je rostoucí, proto x blíže k 0 budeme hodnotit jako 2, blíže k 1 (a na druhé straně blíže k -1) jako 1.

Hledáme t tak, aby:

$$\int_{0 < x < t} x dx = \int_{t < x < 1} (1 - x) dx$$

$$0 = [0,5x^2]_0^t - [x - 0,5x^2]_t^1$$

$$0 = (0,5t^2 - 0) - (1 - 0,5 - t + 0,5t^2) = 0,5t^2 - 0,5 + t - 0,5t^2 = t - 0,5 = 0, \text{ odkud } t = 0,5$$

Meze pro x jsou $t_1 = -0,5$ a $t_2 = 0,5$

Mez pro $r(x)$ je $\theta = p(t_1)/p(t_2) = 0,5/(1 - 0,5) = 1$

Získáváme klasifikátor:

$q(x) = 1$ iff $r(x) > 1$ iff x je z $<-1, -0,5> \cup (0,5, 1>$

$q(x) = 2$ iff $r(x) \leq 1$ iff x je z $<-0,5, 0,5>$

Chyba na třídách $\epsilon_q(1) = \epsilon_q(2) = 2 * \int_{0 < x < 0,5} x dx = 2 * [0,5x^2]_0^{0,5} = 2 * (1/8 - 0) = 1/4$

3.1

a)

Model bude mít jeden parametr $p \in (0, 1)$, $p(H) = p$, $p(T) = 1 - p$, $K = \{H, T\}$, $T = (H, H, H, T, T)$

$$L(p) = p(T|p) = \prod_{i=1}^N p(x_i) = \prod_{k \in K} p(k)^{n_k} = p^{n_H} (1-p)^{n_T}$$

$$l(p) = n_H \ln(p) + n_T \ln(1-p)$$

$$l'(p) = n_H / p + n_T \cdot (-1) / (1-p) = 0$$

Získáme: $n_H / p = n_T / (1-p)$, odkud $1-p = p \cdot n_T / n_H$ a $p = n_H / (n_H + n_T)$,

což pro $n_H = 3$, $n_T = 2$ dává $p = 3/5$

b)

Tady $\eta = \eta$:

$$p(H) = 1 / (1 + e^{-\eta}), \text{ odkud } p(T) = 1 / (1 + e^{\eta})$$

$$L(\eta) = \prod_{k \in K} p(k)^{n_k} = 1 / ((1 + e^{-\eta})^{n_H} \cdot (1 + e^{\eta})^{n_T})$$

$$l(\eta) = -\ln((1 + e^{-\eta})^{n_H} \cdot (1 + e^{\eta})^{n_T}) = -n_H \ln(1 + e^{-\eta}) - n_T \ln(1 + e^{\eta})$$

$$l'(\eta) = -n_H \cdot e^{-\eta} \cdot (-1) / (1 + e^{-\eta}) - n_T \cdot e^{\eta} / (1 + e^{\eta}) = 0$$

$$n_H \cdot e^{-\eta} / (1 + e^{-\eta}) = n_T \cdot e^{\eta} / (1 + e^{\eta})$$

$$n_H / (1 + e^{\eta}) = n_T / (1 + e^{-\eta})$$

Pak přes roznásobení a kvadratickou rovnici $2 \cdot e^{2\eta} - e^{\eta} - 3 = 0$ k $e^{\eta} = 3/2$ a $p(T) = 2/5$

Ta sigmoid funkce je logistická regrese, rozhoduje pomocí lineární funkce poměru $p(1|x)/p(2|x)$. Lze použít pouze u některých problémů a tento (multinomiální naivní Bayes) je jedním z nich, dále normální rozdělení se stejnou kovarianční maticí a nezávislé binární vstupní vektory.

Problém by asi teoreticky nastal při $p(H) = 0$ nebo 1, protože obor hodnot sigmoidy je $(0, 1)$.

3.2

a)

To znamená, že ponožku vrátíme po vytazení zpatky

b)

$$p \cdot p$$

c)

$$l(N=10, R=2) = (10 \text{ nad } 2) \cdot p^2 \cdot (1-p)^8$$

$$l(N, R) = (N \text{ nad } R) \cdot p^R \cdot (1-p)^{N-R}$$

d)

$$L(p) = p^R \cdot (1-p)^{N-R}$$

$$l(p) = R \ln(p) + (N-R) \ln(1-p)$$

$$l'(p) = R/p - (N-R)/(1-p) = 0, \text{ což vede na } p = R/N$$

3.3

a)

$$L(\mu, \Sigma) = \prod_{x \in T} p(x)$$

$$L(\mu, \Sigma) = \prod_{x \in T} (2\pi)^{-d/2} (\det(\Sigma))^{-1/2} e^{-1/2 (x-\mu)^T \Sigma^{-1/2} (x-\mu)}$$

$$l(\mu, \Sigma) = \sum_{x \in T} (-1/2 \ln((2\pi)^d \cdot \det(\Sigma)) - 1/2 (x-\mu)^T \Sigma^{-1/2} (x-\mu))$$

$$l(\mu, \Sigma) = \sum_{x \in T} (-1/2 \ln((2\pi)^d \cdot \det(\Sigma)) - 1/2 x^T \Sigma^{-1/2} x + x^T \Sigma^{-1/2} \mu - 1/2 \mu^T \Sigma^{-1/2} \mu)$$

$$l'(\mu, \Sigma)_\mu = \sum_{x \in T} (x^T \Sigma^{-1/2} - \mu^T \Sigma^{-1/2}) = 0$$

$$\sum_{x \in T} x^T \Sigma^{-1/2} = N \cdot \mu^T \Sigma^{-1/2} \quad (\Sigma^{-1/2} \text{ musí mít inverzi})$$

$$1/N \cdot \sum_{x \in T} x = \mu$$

b) ?????

$$l(\mu, \Sigma) = \text{Sum}_{x \in T} (-\frac{1}{2} \ln((2\pi)^d \det(\Sigma)) - \frac{1}{2} (x - \mu)^T \Sigma^{-1/2} (x - \mu))$$

$$l'(\mu, \Sigma)_{\Sigma} = \text{Sum}_{x \in T} (-\frac{1}{2} \Sigma^{-1} - \frac{1}{2} (x - \mu) (x - \mu)^T) = 0$$

$$\text{Sum}_{x \in T} \Sigma^{-1} = \text{Sum}_{x \in T} (x - \mu) (x - \mu)^T \dots \text{což není správně ale, překáží tam ta inverze u } \Sigma$$

3.4

a)

$$L(h) = p(T|h) = \text{Prod}_{1 \leq i \leq N} h^* e^{-h^* t_i}$$

$$l(h) = \text{Sum}_{1 \leq i \leq N} \ln(h^* e^{-h^* t_i}) = n \ln(h) - \text{Sum}_{1 \leq i \leq N} h^* t_i$$

$$l'(h) = n/h - \text{Sum}_{1 \leq i \leq N} t_i = 0, \text{ odtud } h = n / (\text{Sum}_{1 \leq i \leq N} t_i) = (\text{ar. průměr})^{-1}$$

b)

Expected lifetime je stredni hodnota $t = E(t) = \text{Integral}_{-\infty, \infty} t^* p(t) dt$

$$E(t) = \text{Integral}_{-\infty, \infty} t^* \frac{1}{\theta} e^{-t/\theta} dt = | \text{per partes: } u = t \rightarrow u' = 1, v' = e^{-t/\theta} \rightarrow v = e^{-t/\theta} \cdot -\theta | =$$

$$[-\theta t^* e^{-t/\theta}]_0^{\infty} - \frac{1}{\theta} \text{Integral}_{-\infty, \infty} -\theta e^{-t/\theta} dt = 0 + \text{Integral}_{-\infty, \infty} e^{-t/\theta} dt = [-\theta e^{-t/\theta}]_0^{\infty} = \theta$$

θ říká za jak dlouho se rozbije průměrně žárovka, h říká kolik žárovek se rozbije za jednotku času.

c)

$$h^* = \text{argmax}_h p(h|T)$$

$$p(h|T) = p(T|h) * p(h) = e^{-h} * \text{Prod}_i h^* e^{-h^* t_i}$$

$$\ln(p(T|h) * p(h)) = -h + \text{Sum}_i \ln(h^* e^{-h^* t_i}) = -h + n \ln(h) - h^* \text{Sum}_i t_i$$

$$\text{derivative} = -1 + n/h - \text{Sum}_i t_i = 0, \text{ odtud } h_{\text{MAP}} = n / (1 + \text{Sum}_i t_i)$$

Očekavana životnost $O = \text{ar. průměr} + 1/n$

3.5

a)

$$Z \text{ predchoziho prikladu mame } h = n / (\text{Sum}_i t_i) = 3 / 600 = 1 / 200 = 0.005$$

b) ?????

Nevim jak na to vubec teda

3.6

a)

$$L(\theta) = 1 / \theta^n, \max L(\theta) \text{ pro } \theta \rightarrow 0, \text{ ale } \theta \geq \max_i x_i, \text{ proto } \theta = \max_i x_i$$

b)

Ano, protože jsme pravděpodobně neviděli maximální číslo.

c)

$$\theta \leq M, p(\theta) = 1/M \text{ (uniformní rozdělení)}$$

$$p(\theta|T) = p(T|\theta) * p(\theta) = 1 / (\theta^n * M), \text{ což zas vede na } \theta = \max_i x_i$$

d) ?????

Aposteriorní rozdělení je na $(\max_i x_i, \infty)$, s tím že nižší hodnoty jsou pravděpodobnější

Nevím, jestli chtějí přímo vzorec, pak $p(\theta|T) = p(T|\theta) * p(\theta)$, kde ale $p(\theta)$ nevím jak popsat, možná se dá zanedbat, protože pokud je θ náhodná veličina, každá hodnota má stejnou pravděpodobnost

$$\theta^* = \text{argmin}_{\theta} R(\theta) = \text{argmin}_{\theta} \text{Integral}_{\max x_i, \infty} p(t|T) * (t - \theta)^2 dt$$

Podle wikipedie

(https://en.wikipedia.org/wiki/Bayes_estimator#Minimum_mean_square_error_estimation) by

$$\text{to mělo být } \theta^* = \text{Integral}_{\max x_i, \infty} \theta^* p(\theta|T) dt = \text{Integral}_{\max x_i, \infty} \theta / \theta^n dt = \text{Integral}_{\max x_i, \infty} \theta^{-n+1} dt =$$

$$[1/(-n+2) * \theta^{-n+2}]_{\max x_i}^{\infty} = \infty, \text{ což je asi kravina (bral jsem zde ten zanedbaný } p(\theta), \text{ viz výše)}$$

4.1

$$L(d) = p(T|d) = \prod_{1 \leq i \leq K} d_i^{n_i}$$

$$l(d) = \sum_{1 \leq i \leq K} n_i \ln(d_i), \sum_{1 \leq i \leq K} d_i / K = 1$$

$$\text{Lagrange} = \sum_{1 \leq i \leq K} n_i \ln(d_i) + \lambda' * (\sum_{1 \leq i \leq K} d_i - K)$$

Derivace Lagrange podle $d_i = n_i/d_i + \lambda' = 0$, odtud $n_i/d_i = \lambda = \text{konst}$, $d_i = n_i/\lambda$

Víme $\sum_{1 \leq i \leq K} d_i = K$, odtud $\sum_{1 \leq i \leq K} n_i / \lambda = K$, což je $N / \lambda = K$

Máme $\lambda = N / K$, takže $d_i = n_i/\lambda = K * n_i/N$

4.2

Máme trénovací data $T = \{(x_i, k_i)\}$. Pro testovanou hodnotu x nalezneme K nejbližších x_j z T . Pro x volíme třídu k takovou, která má mezi x_j většinu.

Plusy:

- jednoduchost (při naivní implementaci)

Mínusy:

- přefitovává se
- časová náročnost (při naivní implementaci prochází trénovací body lineárně, lze zlepšit použitím k-d stromů a/nebo pravděpodobnostními zárukami nejbližšího souseda)
- paměťová náročnost (při naivní implementaci potřebuje všechny trénovací body, jde zlepšit redukcí trénovací množiny)

1-NN: nejbližší je (4.5, A), volíme A

3-NN: nejbližší jsou (4.5, A), (6, B) a (3, B), volíme B

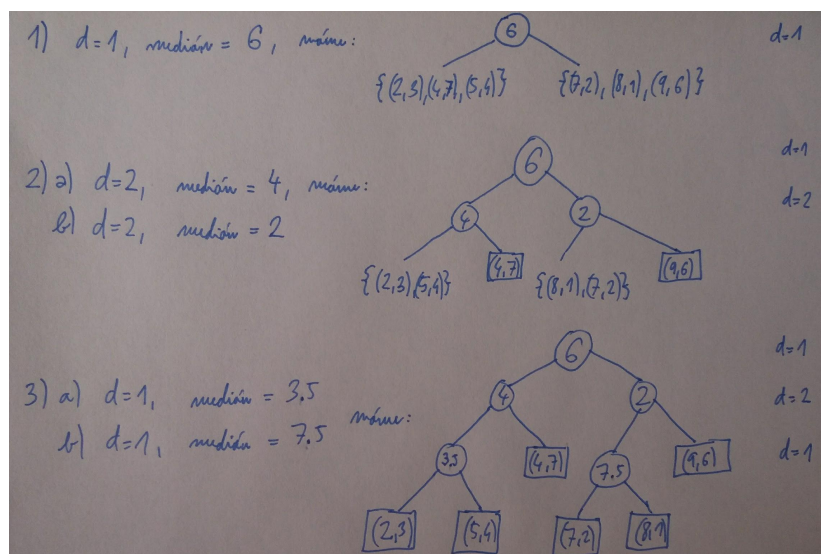
5-NN: nejbližší jsou (4.5, A), (6, B) a (3, B), (2, A) a (1.5, A), volíme A

4.3

Algoritmus:

- 1) Začínáme na indexu vstupních vektorů $d = 1$ (beru v matici one-indexed), aktuální data = trénovací množina
- 2) Pokud počet aktuálních dat je 1, ukončí větev rekurze.
- 3) Najdeme medián aktuální části dat podle d -tého indexu a rozdělíme data na dvě části a ty dáme do větví aktuálního uzlu
- 4) Rekurzivně pokračujeme do větví krokem 2) a s indexem $d+1$ (modulo D)

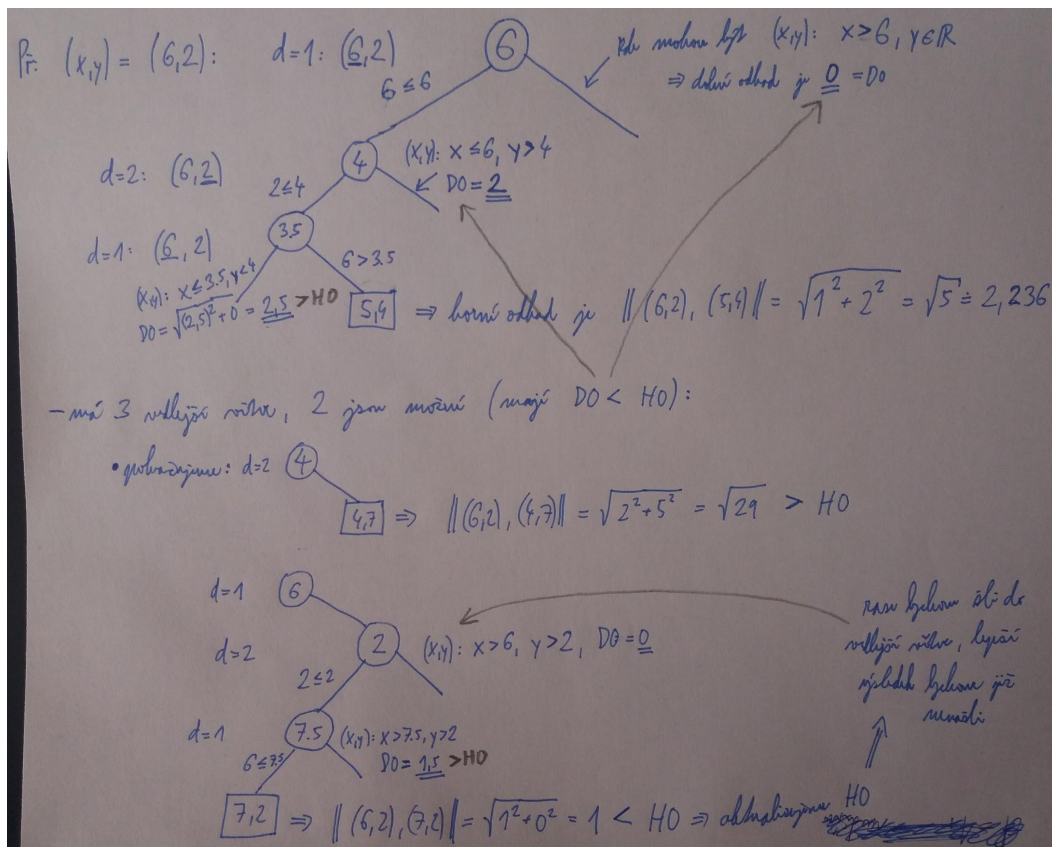
Vytvoření k-D stromu:



Hledání NN:

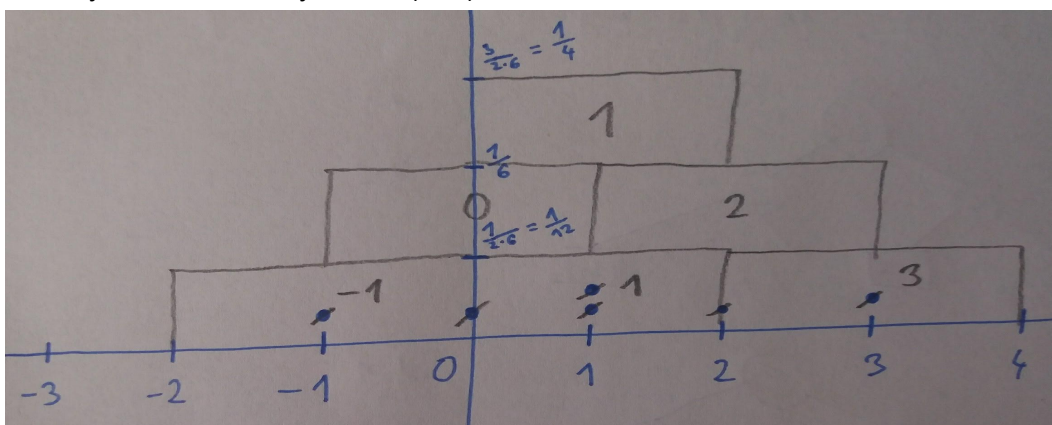
Máme bod (x, y) . Procházíme od kořene dolů, vždy se vydáváme do té větve, kam nás zavede hodnota x nebo y podle d a větvičí nerovnosti v dané hloubce. Při každém zanoření se nám u druhé (nenavštívené) větve určí nejbližší možný soused v ní. Po doražení do listu máme horní odhad nejbližšího souseda. Navštívíme ty větve, kde mohou být bližší sousedi.

Příklad (k-D ušetřilo porovnávání s (2, 3) a (8, 1), což není moc teď zrovna, ale při větších množinách to samozřejmě je v průměru více znát):



4.4

Integrál pod kombinací všech jádrových funkcí musí být 1, proto do každého bodu nevložíme obdélník s výškou $\frac{1}{2}$, ale s výškou $1/(2 \cdot N) = 1/12$:



4.5 ?????

Máme Lagrange = $\sum_i \sum_j K_{ij}/K_i * \log(\pi_j * K_i) + h' * (\sum_i \pi_i - 1)$

“Derivace Lagrange podle π_j ” = $\sum_i K_{ij}/(K_i * \pi_j) + h' = 0$

Odtud $\pi_j = 1/h * \sum_i K_{ij}/K_i$ ($h = -h'$)

Po dosazení do $\sum_i \pi_i = 1$ máme $h = \sum_i \sum_j K_{ij}/K_i$

5.1

a)

Pres $p(1|x)/p(-1|x) = e^{w^T x}$ získáme $p(1|x) = 1 / (1 + e^{-w^T x})$, $p(-1|x) = 1 / (1 + e^{w^T x})$

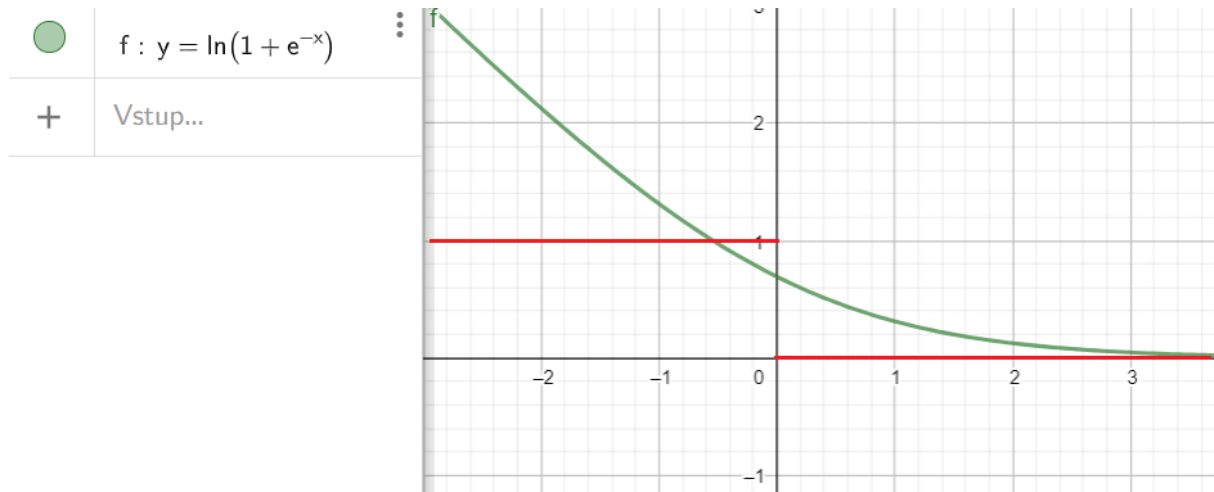
b)

$p(x, k) = p(k|x) * p(x) \rightarrow \ln(p(x, k)) = \ln(p(k|x)) + \ln(p(x))$, kde $\ln(p(x))$ je const

$p(k|x) = 1/(1 + e^{-k^T w^T x})$, tedy negace součtu přes všechna trénovací data $E(w)$:

$$E(w) = -l(w) = -\sum_i \ln(1/(1 + e^{-k_i^T w^T x_i})) = \sum_i \ln(1 + e^{-k_i^T w^T x_i})$$

c)



$E(w)$ je součet konvexních funkcí $\Rightarrow$ je konvexní

5.2

a)

$$1 - 1/(1 + e^{-z}) = (1 + e^{-z} - 1)/(1 + e^{-z}) = e^{-z}/(1 + e^{-z}) = 1/(1 + e^z)$$

b)

$$\text{Derivace } 1/(1 + e^{-z}) = -1 * e^{-z} * (-1)/(1 + e^{-z})^2 = e^{-z}/(1 + e^{-z})^2 = e^{-z}/(1 + e^{-z}) * 1/(1 + e^{-z}) = 1/(1 + e^z) * 1/(1 + e^{-z})$$

c)

$$\text{Derivace } \ln(\sigma(z)) = 1/\sigma(z) * (\text{Derivace } \sigma(z)) = 1/\sigma(z) * \sigma(z) * (1 - \sigma(z)) = 1 - \sigma(z)$$

d)

$$\begin{aligned} \text{Derivace } \ln(\sigma(w^T x)) \text{ podle } x &= (\text{Derivace } \ln(\sigma(w^T x)) \text{ podle } w^T x) * (\text{Derivace } w^T x \text{ podle } x) = \\ &= (1 - \sigma(w^T x)) * w^T \end{aligned}$$

e)

Derivace $-\ln(\sigma(z)) = -\sigma(-z) = \sigma(z) - 1$, což je rostoucí funkce

Druhá derivace $-\ln(\sigma(z)) = \text{Derivace } -\sigma(-z) \text{ podle } z = -1 * (\text{Derivace } \sigma(-z) \text{ podle } -z) *$

$(\text{Derivace } -z \text{ podle } z) = -1 * (\text{Derivace } \sigma(z) \text{ podle } z) * -1 = \sigma(z)(1 - \sigma(z))$, což je součin

kladných čísel, který je kladný

Problem 6.1

a)

Lagrangeova funkce:

$L = \|x-y\|^2 + \lambda(a^T y + b)$, x , a , b jsou konst, y proměnná a λ multiplikátor

Podmínka: $a^T y + b = 0$

Derivace L podle $y = -2(x-y) + \lambda a = 0$, odtud:

$$2(x-y) = \lambda a$$

[pokračování od Káji Balagazové]

Dosadíme $y = x - \lambda a/2$ do $a^T y + b = 0$ a dostaneme λ bdu:

$$\lambda = 2(a^T x + b)/(a^T a)$$

a výsledná vzdálenost je:

$$\|x-y\| = \|x - x + \lambda a/2\| = |\lambda/2| \|a\| = (|a^T x + b| / \|a\|^2) * \|a\| = |a^T x + b| / \|a\|$$

[Od Kačky Mackové:]

Jde nad tím přemýšlet i z úhlu pohledu lineární algebry. Rozhodovací nadrovina je dána normálovým vektorem a a posunem b . Hledaná vzdálenost je délka projekce bodu x na tento normálový vektor posunutý o b . Pak není třeba počítat Lagrange. Vektor a vydělíme $\|a\|$, aby byl jednotkový. Promítneme na něj bod x :

$\text{proj}(x) = a^* (a^T x + b) / \|a\|^2$ a máme projekci.

Projekci znormujeme a máme:

$$\|\text{proj}(x)\| = \|a\| * |a^T x + b| / \|a\|^2 = |a^T x + b| / \|a\| = \|x-y\|$$

b)

Máme $x \in \mathbb{R}^D$ a $k \in \{1, -1\}$. Definujeme $x' = k_i^*(x_i, 1) \in \mathbb{R}^{D+1}$ a $w = (a, b) \in \mathbb{R}^{D+1}$

Pak $w^T x' = k_i(a^T x_i + b) \geq 0$ iff $\text{sign}(a^T x_i + b) = k_i$ tj. iff $q(x) = k_i$ neboli q klasifikuje správně

Problem 6.2

a)

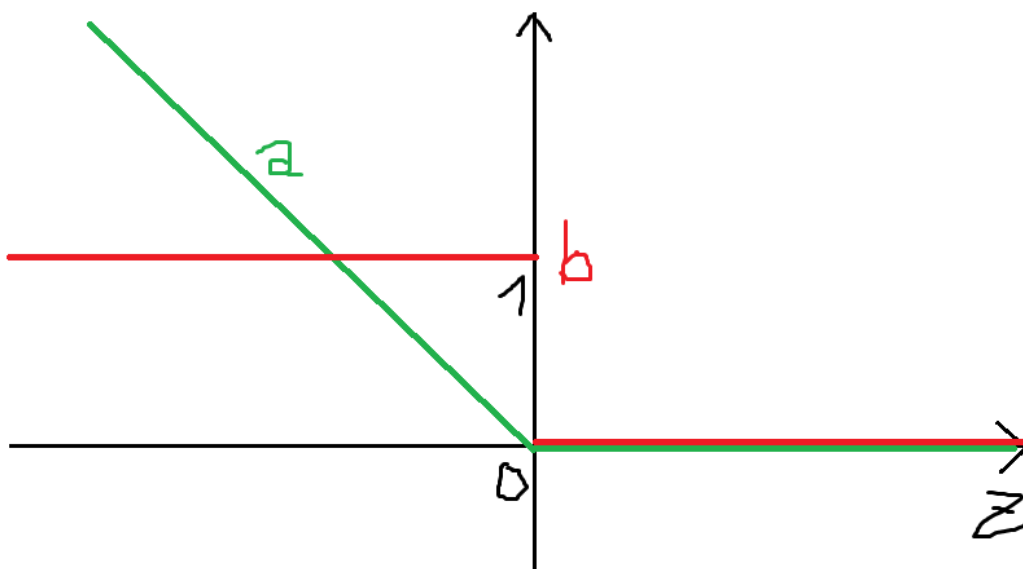
$\max(-z, 0)$ (zeleně)

b)

$R_{\text{emp}}(q(x)) = W(q(x), k) = 1$ iff $q(x) \neq k$ (tj. iff $z < 0$), 0 else (červeně)

c)

logaritmus ze cvičení 5.1c)



Problem 6.3

a)

Přepíšeme si:

$$l_i(w) = \begin{cases} -w^T x_i & \text{iff } x_i \text{ je špatně ohodnocený} \\ 0 & \text{iff } x_i \text{ je správně ohodnocený} \end{cases}$$
$$\nabla_w l_i(w) = \begin{cases} -x_i & \text{iff } x_i \text{ je špatně ohodnocený} \\ 0 & \text{iff } x_i \text{ je správně ohodnocený} \end{cases}$$

Takže správně ohodnocený bod $w^{t+1} = w^t + 0$ nezmění a špatně ohodnocený updatuje $w^{t+1} = w^t + \varepsilon^* x_i$

b)

Intuice říká, že na velikosti výsledného vektoru w nezáleží, jde pouze o jeho směr. Výsledný vektor $w = \varepsilon^*(\text{lineární kombinace vektorů } x)$, po vydělení ε je směr stále stejný. Násobením ε ($\varepsilon > 0$) se sice změní velikost $w^T x_i$, ale znaménko zůstane stejné (a znaménko je to, co nás zajímá).

Máme $w^0 = 0$, hledáme x_i : $w^{0T} x_i \leq 0$. To je hned pro první x_1 .

Updatujeme $w^1 = 0 + \varepsilon^* x_1 = \varepsilon^* x_1$. Hledáme x_i : $w^{1T} x_i = \varepsilon^* x_1^T x_i \leq 0$. ε je ale kladná konstanta, proto lze nerovnost vydělit bez změny znaménka, neboli na ε nezáleží.

Problem 6.4

Upravíme si T na $T' = \{(2, -1, -1), (0, 0, -1), (0, 2, 1), (0, 3, -1), (2, 2, 1)\}$

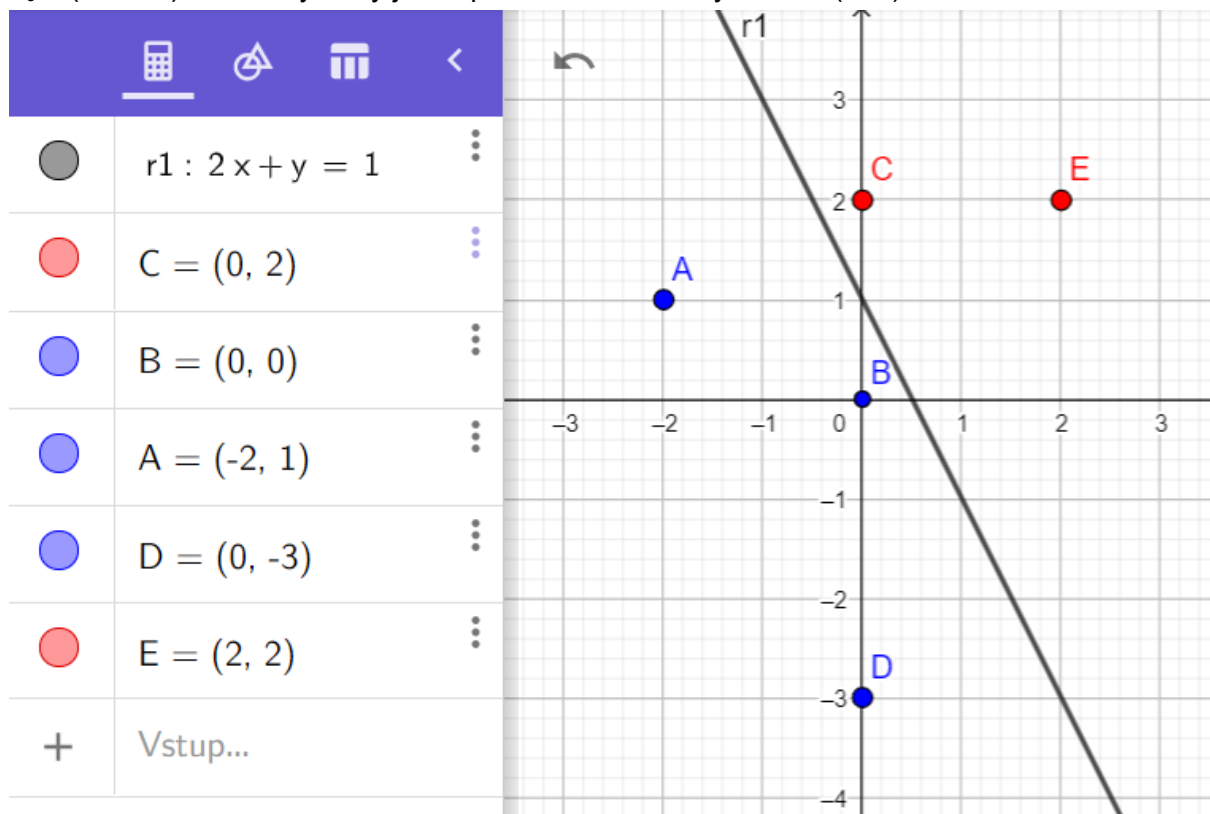
Pracujeme pouze se skalárním součinem $w^T x$ (bez b), chceme $w^T x \geq 0$ pro všechny body, inicializujeme $w_0 = (0, 0, 0)$ a sekvenčně procházíme množinu T .

$w_0 = (0, 0, 0)$ - chyba pro bod A

$w_1 = (2, -1, -1)$ - chyba pro bod C

$w_2 = (2, 1, 0)$ - chyba pro bod B

$w_3 = (2, 1, -1)$ - všechny body jsou správně ohodnoceny $\Rightarrow w = (2, 1)$, $b = -1$



Problem 6.5

Základní krok ($t = 0$):

$$\|w^0\|^2 = 0 \leq 0 \cdot \max \|x_i\|^2 = 0$$

Indukční krok:

Indukční předpoklad: $\|w^t\|^2 \leq t \cdot \max \|x_i\|^2$ (chceme dokázat platnost pro $t+1$)

$$\|w^{t+1}\|^2 = \|w^t + x^t\|^2 = (w^t + x^t)^T (w^t + x^t) = w^t w^t + 2 \cdot w^t x^t + x^t x^t = \dots$$

[pokračování od Káji Balagazové]

$$\begin{aligned} \dots &= \|w^t\|^2 + 2w^t x^t + \|x^t\|^2 \stackrel{\leq_1}{\leq} \|w^t\|^2 + \|x^t\|^2 \stackrel{\leq_2}{\leq} \|w^t\|^2 + \max \|x\|^2 \stackrel{\leq_3}{\leq} t \cdot \max \|x\|^2 + \max \|x\|^2 \stackrel{=}{=} \\ &= (t + 1) \cdot \max \|x\|^2 \end{aligned}$$

$\leq_1$.. x^t je špatně klasifikovaný vektor, tzn $w^t x^t \leq 0$, takže se můžeme zbavit $2 \cdot w^t x^t$

$\leq_2$.. protože $\|x_i\|^2 \leq \max_i \|x_i\|^2$

$\leq_3$.. podle indukčního předpokladu $\|w^t\|^2 \leq t \cdot \max \|x_i\|^2$

$=_4$.. a stačí jenom vytknout $\max \|x\|^2$ a máme dokázáno, že $\|w^{t+1}\|^2 \leq (t + 1) \cdot \max \|x\|^2$

Problem 7.1

a)

Řešení je jednoduché zamyšlení:

Minimalizujeme výraz, kde je ξ v kladném násobku, neboli ho chceme pro minimalizaci co nejmenší. Máme ale omezení $\xi \geq 0$ a $(w^T x + b)y \geq 1 - \xi$. Zvolíme tedy nejmenší takové $\xi \geq 0$, aby $(w^T x + b)y \geq 1 - \xi$ platilo.

Formálněji:

$\min_{w,b,e} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i = \min_{w,b,e} F(w,b,e)$ za podmínky $(w^T x_i + b)y_i \geq 1 - \xi_i$, $\xi_i \geq 0$

Derivace F podle $\xi_i = C$, kde C je kladná konstanta určující ztrátu za špatně ohodnocené x_i

Pro minimalizaci volíme nejmenší ξ_i , které splňuje podmínky.

$\xi_i = 0$ iff $(w^T x_i + b)y_i \geq 1$ (tj. iff x_i je správně ohodnoceno)

$\xi_i = 1 - (w^T x_i + b)y_i$ iff $(w^T x_i + b)y_i < 1$ (tj. iff x_i je špatně ohodnoceno)

Lze to případně řešit i opravdu graficky jako průnik polorovin $\xi_i \geq 0$ a $\xi_i \geq 1 - (w^T x_i + b)y_i$.

Na těchto polorovinách hledáme nejmenší hodnotu funkce $K + C \sum_i \xi_i$, kde K je konstantní součet zbytku původní funkce $K = \frac{1}{2} \|w\|^2 + C \sum_{j \neq i} \xi_j$.

b)

- Přejít od vzorce a) do b):

Máme $\min_{w,b,e} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i$. C je konstanta, hledáme argmin, hodnota min nás nezajímá, proto může vydělit. Získáme:

$\min_{w,b,e} \frac{1}{2C} \|w\|^2 + \sum_i \xi_i$

Za ξ_i dosadíme nejmenší hodnotu tak, aby platily podmínky $(w^T x_i + b)y_i \geq 1 - \xi_i$, $\xi_i \geq 0$ (viz a)). Tahle hodnota se dá kompaktněji napsat jako $\max(1 - (w^T x_i + b)y_i, 0)$. Tím se zbavíme proměnné ξ_i . Získáme:

$\min_{w,b,e} \frac{1}{2C} \|w\|^2 + \sum_i \max(1 - (w^T x_i + b)y_i, 0)$

- Přejít od vzorce a) do b) [od Vítka Lupinka]:

Taky se dá hezky vyřešit úvahou/graficky jako a) máme tam $\xi \geq 0$ a $(w^T x + b)y \geq 1 - \xi$, což když vyplotíme, tak je vlastně obrázek c) a z něj je hezky vidět, že $\xi_i = \max(1 - (w^T x_i + b)y_i, 0)$ tudíž to můžu udělat pro všechny a dostanu: $\min_{w,b} \frac{1}{2C} \|w\|^2 + \sum_i (\max(1 - (w^T x_i + b)y_i, 0))$.

- Hledání minima:

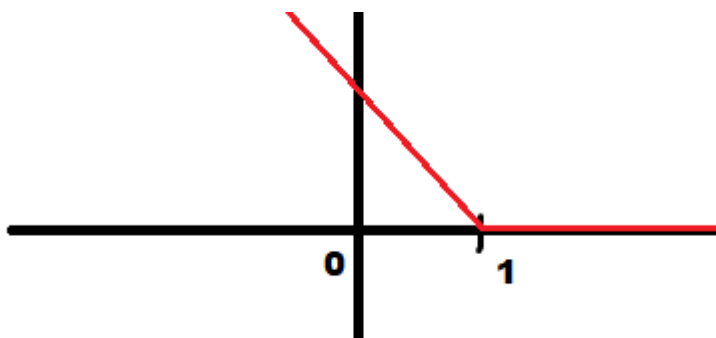
Řekl bych, že si pro každé w a b můžeme vypočítat gradient a použít tedy iterační metody optimalizace:

$F = \min \frac{1}{2C} \|w\|^2 + \sum (\text{přes špatně ohodnocené } i) (1 - (w^T x_i + b)y_i)$

F derivace podle $w = w/C - \sum (\text{přes špatně ohodnocené } i) y_i x_i = 0$

F derivace podle $b = \sum (\text{přes špatně ohodnocené } i) -y_i = 0$

c)



Z tohoto plyne, že i body na správné straně od rozdělovací nadroviny, ale blíže než 1, jsou považovány za špatně ohodnocené. U Perceptronu jsou za špatně ohodnocené považovány opravdu pouze z = $(w^T x_i + b)y_i < 0$. U logistické regrese je nenulová ztráta pro všechny z .

d) ?????

Podle mě ne, když $C \rightarrow \infty$, tak vypadne první člen a máme

$$\min \sum_i \max(1 - (w^T x_i + b)y_i, 0),$$

což má smysl u neseparabilních dat, chceme najít oddělující přímku, tak aby byly všichni správně ohodnoceni nebo když už špatně, tak co nejbliž.

Perceptron v takové situaci nic neříká, nefunguje.

e)

IF Vybraný bod je správně ohodnocený:

$$F = \|w\|^2 / (2nC) + 0$$

$$F'_w = w / (nC)$$

$$F'_b = 0$$

IF Vybraný bod je špatně ohodnocený:

$$F = \|w\|^2 / (2nC) + 1 - (w^T x_i + b)y_i$$

$$F'_w = w / (nC) - x_i y_i$$

$$F'_b = -y_i$$

Funkce roste nejstrměji ve směru gradientu, jdeme v opačném směru s velikostí kroku ϵ :

$$w_{t+1} = w_t - \epsilon F'_w$$

$$b_{t+1} = b_t - \epsilon F'_b$$

Problem 7.2

a)

[řešení od Davida Krause]

Víme $n = w / \|w\|$, $w^T x_0 + b = 0$ neboli $w^T x_0 = -b$

$$n^T (x - x_0) = w^T (x - x_0) / \|w\| = (w^T x - w^T x_0) / \|w\| = (w^T x + b) / \|w\|$$

To je vzdálenost bodu od přímky (na jednu stranu plus, na druhou minus). Když to celé vynásobím y (± 1) a bod bude správně klasifikován, výsledek bude kladný.

b)

$w^*, b^* = \operatorname{argmax}_{w, b} \min_{x, y} (w^T x + b)y / \|w\|$ za podmínky $(w^T x + b)y / \|w\| \geq 0$ (dělení $\|w\|$ v podmínce není nutné, protože $\|w\|$ nezmění znaménko.

c)

- To je tak jasný, až nevím jak to dokazovat. Asi takhle:

1) $\min_i d_i \leq d_{\max}$: kdyby $\min_i d_i > d_{\max}$, tak lze $d_{\max}$ zvětšit, tedy $d_{\max}$ není $\max(d: d \leq d_i)$

2) $\min_i d_i \geq d_{\max}$: kdyby $\min_i d_i < d_{\max}$, tak neplatí podmínka v $\max(d: d \leq d_i)$

Proto $\min_i d_i = d_{\max}$.

- Přejít k nové formulaci:

$w^*, b^* = \operatorname{argmax}_{w, b} \min_{x, y} (w^T x + b)y / \|w\|$ za podmínky $(w^T x + b)y / \|w\| \geq 0$ pro $\forall (x, y)$

Přidáme novou proměnnou d :

$\operatorname{argmax}_{w, b, d} d$ za podmínky $(w^T x + b)y / \|w\| \geq 0$ pro $\forall (x, y)$ a $d \leq \min_{x, y} (w^T x + b)y / \|w\|$

$\operatorname{argmax}_{w, b, d} d$ za podmínky $(w^T x + b)y / \|w\| \geq 0$ a $d \leq (w^T x + b)y / \|w\|$ pro $\forall (x, y)$

$\operatorname{argmax}_{w, b, d} d$ za podmínky $0 \leq d \leq (w^T x + b)y / \|w\|$ pro $\forall (x, y)$

d)

No jasný, zajímá nás znaménko výrazu $w^T x + b$, které se nezmění vynásobením kladnou konstantou. Přesněji:

$(w^T x + b)y / \|w\| \geq d \geq 0$, substituujeme:

$$(s^* w^T x + s^* b)y / \sqrt{(s^* w)^T s^* w} \geq d \geq 0$$

$$s^* (w^T x + b)y / |s^*| \|w\| \geq d \geq 0 \quad (s > 0)$$

$$(w^T x + b)y / \|w\| \geq d \geq 0$$

e)

$\max_{w,b} 1 / \|w\|$ za podmínky $(w^T x + b)y \geq 1$ pro $\forall (x, y)$

f)

Protože $\|w\| \geq 0$, tak $\arg\max 1/\|w\|$ je pro $\arg\min \|w\|$ a $\arg\min \|w\|^2$ je pro $\arg\min \|w\|$

Problem 7.3

a)

$\min_{w,b,\xi} \max_{\alpha} \frac{1}{2} w^T w + C 1^T \xi - (w^T X + by^T - 1^T + \xi^T) \alpha$ za podmínky $\xi \geq 0, \alpha \geq 0$

Význam:

Pokud jsou všechny body správně ohodnoceny (tj. $w^T X + by^T - 1^T + \xi^T \geq 0$), tak tak $-(w^T X + by^T - 1^T + \xi^T) = -G$ je záporné, my chceme $\max -G^* \alpha$, kde jsou v tu chvíli α násobeny zápornými koeficienty, proto je max pro minimální α , neboli pro $\alpha = 0$, a tedy $-G^* \alpha = 0$. Pokud je nějaký bod x_i špatně klasifikován, $-G_i^* \alpha > 0$ a tedy lze maximalizovat maximalizací $\alpha_i = \infty$, proto $-G^* \alpha = \infty$

b) ?????

Obecně platí: $\max_a \min_b \leq \min_b \max_a$

Máme: $\max_{\alpha} \min_{w,b,\xi} \frac{1}{2} w^T w + C 1^T \xi - (w^T X + by^T - 1^T + \xi^T) \alpha$ za podmínky $\xi \geq 0, \alpha \geq 0$

Řešíme vnitřní $\min_{w,b,\xi} f(w,b,\xi) = \min_{w,b,\xi} \frac{1}{2} w^T w + C 1^T \xi - w^T X \alpha - by^T \alpha + 1^T \alpha - \xi^T \alpha$

w:

Transponujeme a pak derivace f podle w = $w^T - \alpha^T X^T = 0$ a odtud $w = X \alpha$

ξ :

Dáme k sobě členy $C 1^T \xi - \xi^T \alpha = C 1^T \xi - \alpha^T \xi = (C 1 - \alpha)^T \xi$

$\min_{\xi \geq 0} (C 1 - \alpha)^T \xi$, kde $(C 1 - \alpha)$ jsou vlastně koeficienty pro ξ

Pokud $(C 1 - \alpha) \geq 0$, minimalizujeme výraz minimální ξ , neboli $\xi = 0$, tak $(C 1 - \alpha)^T \xi = 0$

Pokud existuje $(C - \alpha_i) < 0$, lze dosáhnout $-\infty$ tím, že $\xi_i \rightarrow \infty$

Celkově chceme ale max min f, proto nechceme $-\infty$, neboli nám to dává omezení $\alpha \leq C$

b:

Protože b není omezené, pokud jeho koeficient $y^T \alpha$ není nula, lze člen $y^T \alpha b$ minimalizovat volbou $b = +$ nebo $-\infty$ k hodnotě $-\infty$. Zase ale chceme celkově max min, proto získáváme podmínku $y^T \alpha = 0$

Tedy v $\max_{\alpha} \min_{w,b,\xi} \frac{1}{2} w^T w + C 1^T \xi - w^T X \alpha - by^T \alpha + 1^T \alpha - \xi^T \alpha$ za podmínky $\xi \geq 0, \alpha \geq 0$

nahradíme $w = X \alpha$, místo $C 1^T \xi - \xi^T \alpha$ přidáme podmínku $\alpha \leq C$ a místo $by^T \alpha$ přidáme podmínku $y^T \alpha = 0$. Získáme:

$\max_{\alpha} \frac{1}{2} \alpha^T X^T X \alpha - \alpha^T X^T X \alpha + 1^T \alpha$ za podmínky $\alpha \geq 0, \alpha \leq C, y^T \alpha = 0$, neboli

$\max_{\alpha} -\frac{1}{2} \alpha^T X^T X \alpha + 1^T \alpha$ za podmínky $0 \leq \alpha \leq C, y^T \alpha = 0$

Problem 7.4

a)

Vliv mají ty X_i , kde $\alpha_i > 0$

b)

Když je někde ostrá nerovnost, existuje dostatečně malé delta, aby ostrá nerovnost platila i po změně, to je jasné.

Vztah s α asi:

V primární úloze bod (x, y) ostře splňuje $y(w^T x + b) > 1$. Duální úloha je

$\arg\max_a \min_{w,b} (\frac{1}{2} \|w\|^2 - \sum_i a_i (y_i (w^T x_i + b) - 1))$, kde pro náš bod je $y(w^T x + b) - 1 = D > 0$.

Příslušná alfa a tedy je odčítána jako $-D^* a$. Chceme $\arg\max_a$, proto je tato alfa $a = 0$.

Problem 7.5

a)

$\max_a \sum_i (a) - \frac{1}{2} \sum_{i,j} (a_i a_j y_i y_j \phi(x_i) \phi(x_j))$ za podmínky $\sum_i a_i y_i = 0$, $0 \leq a \leq C$

b)

To asi ne, pokud třeba uvažujeme nějaké crazy $K(x, y) = \sin(x_1^5 y_1^5 + \sqrt{x_1 y_1})^{x_2 y_2}$, tak asi nějaký předpis $\phi(x)\phi(y)$ nenajdeme. Může nás spasit možná Taylorův rozvoj, ale ten má odchylku nebo to je nekonečná posloupnost. Kernel funkce může být prostě každá pozitivně semidefinitní funkce $R^D \times R^D \rightarrow R$.

c)

Pro bod z máme $q(z) = \phi(w)^T \phi(z) + b = \sum_{i \in \text{sv}} a_i y_i \phi(x_i)^T \phi(z) + b = \sum_{i \in \text{sv}} a_i y_i K(x_i, z) + b$, kde b předpočítáme (sv = support vector), protože víme:

$y_{\text{sv}}(\phi(w)^T \phi(x_{\text{sv}}) + b) = y_{\text{sv}}(\sum_{i \in \text{sv}} a_i y_i K(x_i, x_{\text{sv}}) + b) = 1$ pro nějaký support vector x_{sv} , a odtud máme:

$\sum_{i \in \text{sv}} a_i y_i K(x_i, x_{\text{sv}}) + b = y_{\text{sv}}$ (protože $y_{\text{sv}} = \pm 1$)

$b = y_{\text{sv}} - \sum_{i \in \text{sv}} a_i y_i K(x_i, x_{\text{sv}})$

Problem 7.6

a)

$K(x, y) = 1 + x_1 y_1 + x_2 y_2 + x_1^2 y_1^2 + x_2^2 y_2^2 + x_1 x_2 y_1 y_2$

b)

$K(x, y) = (1 + x_1 y_1 + x_2 y_2)^2 = 1 + 2x_1 y_1 + 2x_2 y_2 + 2x_1 x_2 y_1 y_2 + x_1^2 y_1^2 + x_2^2 y_2^2$

Odtud máme $\phi(x) = (1, \sqrt{2}x_1, \sqrt{2}x_2, \sqrt{2}x_1 x_2, x_1^2, x_2^2)$

c)

ano, šlo by to, získáme polynom 2D proměnných řádu 2d, mapování bude do dimenze k , kde k = počet členů polynomu, každá dimenze mapování bude obsahovat odmocninu konstanty * součin všech prvků jednoho vektoru v v daném členu

Výpočet kernel function (získání hodnoty) je cca D násobení ($x_1 * y_1, \dots, x_D * y_D$) + D sčítání ($1 + x_1 * y_1 + \dots + x_D * y_D$) + d násobení (mocnění na d).

Pro získání mapování potřebujeme získat předpis kernel function:

D násobení a získáme $(1 + x_1 * y_1 + \dots + x_D * y_D)$, což má $D+1$ členů. Při umocnění na druhou uděláme $(D+1)(D+1)$ násobení a vznikne člen s $D(D+1)$ členy (myslím), takže další násobení původním $(1 + x_1 * y_1 + \dots + x_D * y_D)$ (tj., celkově umocnění na třetí) je $D(D+1)(D+1)$ násobení a vznikne výraz s $D * D * (D+1)$ členy atd. Takže poslední krok mocnění máme asi

$D^{d-1}(D+1)(D+1)$ násobení plus (nepočítal jsem sčítání).

Výpočet mapování pak je odmocnění konstant u členů ve výsledném výrazu s $D^d(D+1)$ členy a rozdělení částí x a y v každém členu.

Nechce se mi nad počty operací přemýšlet nějak do hloubky, prostě kvůli tomu že pro získání mapování potřebujeme předpis a tak nemůžeme po každém kroku členy sečíst do jednoho skaláru, nám přibude strašně moc operací. Asi :D

Problem 8.1

a)

chyba je když $\text{sign}(f(x)) \neq y$, což je iff

$$f(x) < 0 \text{ a } y = 1$$

$$f(x) > 0 \text{ a } y = -1$$

$$\Rightarrow f(x) \cdot y < 0$$

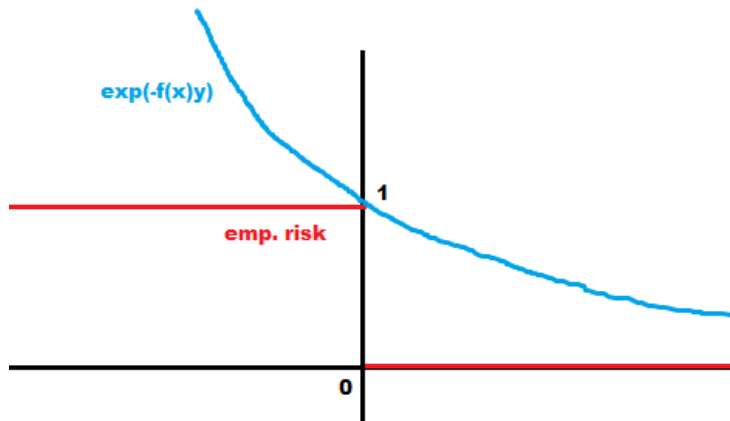
Ty divný závorky dávají 1 iff výraz uvnitř je true, jinak 0, takže to sečte všechny chyby.

Přijde mi jen divný, že to není dělený velikostí trénovací množiny, v přednášce to bylo násobené $D_i(i)$, což bylo inicializované na $1/N$.

b)

Když $f(x)y \leq 0$ tak emp. risk je 1, ale $-f(x)y \geq 0$, takže $e^{-f(x)y} \geq 1$

Když $f(x)y > 0$ tak emp. risk je 0, ale $e^{-f(x)y} > 0$ platí vždy



Výhodou může být, že i správně ohodnocené body mají přidělenou nenulovou chybu, takže se je algoritmus stále snaží ohodnotit "lépe", neboli dostat dál od oddělovací hranice. Zároveň je to spojitá, diferencovatelná funkce.

c)

Máme mapování $h(x): \mathbb{R}^D \rightarrow \mathbb{R}^T$ a chceme najít lineární klasifikátor takto mapovaných bodů.

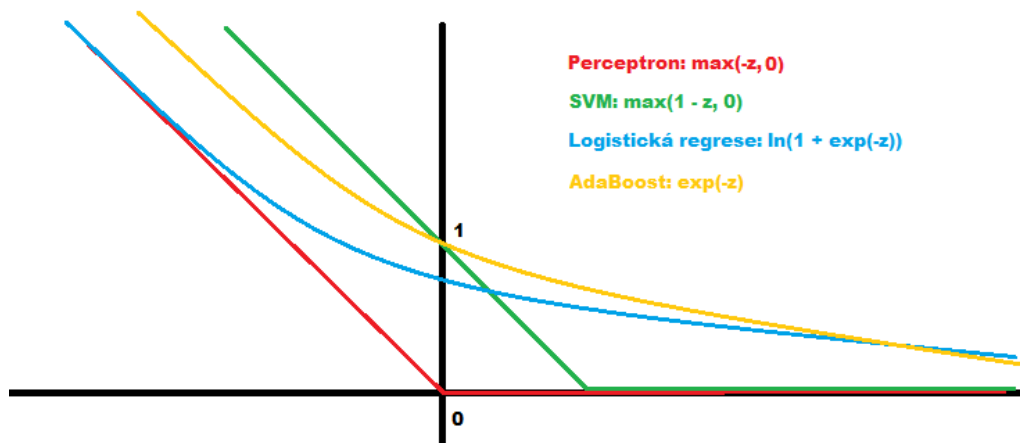
Tzn. hledáme koeficienty a z $\mathbb{R}^T$, tak aby chyba na trénovacích datech $\{(x_i, y_i)\}_{1 \leq i \leq N}$

$R_{\text{emp}} = \sum_{1 \leq i \leq N} e^{-a^T h(x_i) y_i}$ byla minimální $\Rightarrow \text{argmin}_a R_{\text{emp}}$

SVM: $R_{\text{emp}} = \|w\|^2 / (2C) + \sum_{1 \leq i \leq N} \max(1 - (w^T x_i + b) y_i, 0) \Rightarrow \text{argmin}_{w,b} R_{\text{emp}}$

Perceptron: $R_{\text{emp}} = 1/N \sum_{1 \leq i \leq N} \max(-w^T x_i, 0)$

LogReg: $R_{\text{emp}} = \ln(1 + e^{-k_i^T w^T x_i})$



Problem 8.2

a) ?????

$$R = \sum_{1 \leq i \leq N} e^{-\sum_{1 \leq j \leq t} a_j h_j(x_i) y_i - a^* h_t(x_i) y_i} = \sum_i D_t(i) * e^{-a^* h_t(x_i) y_i}$$
$$e^{-\sum_{1 \leq j \leq t} a_j h_j(x_i) y_i} = D_t(i) > 0, \text{ ale nevím jak dokázat to sčítání do 1}$$

Problem 8.3

a)

$$\min_{a, h_t} \sum_{1 \leq i \leq N} D_t(i) * e^{-a^* h_t(x_i) y_i} = \min_{a, h_t} \sum_{i: h_t(x) \neq y_i} D_t(i) * e^{a^*} + \sum_{i: h_t(x) = y_i} D_t(i) * e^{-a^*}$$

b)

$$\begin{aligned} & \sum_{\text{špatně } i} D_t(i) * e^{a^*} + \sum_{\text{dobře } i} D_t(i) * e^{-a^*} = \\ & = \sum_{\text{špatně } i} D_t(i) * e^{a^*} + \sum_{\text{všechny } i} D_t(i) * e^{-a^*} - \sum_{\text{špatně } i} D_t(i) * e^{-a^*} = \\ & = e^{-a^*} * \sum_{\text{všechny } i} D_t(i) + (e^{a^*} - e^{-a^*}) * \sum_{\text{špatně } i} D_t(i) \end{aligned}$$

c)

Pro $a > 0$ je $e^{a^*} - e^{-a^*} > 0$. Předchozí vyjádření závisí na h_t pouze ve členu $(e^{a^*} - e^{-a^*}) * \sum_{\text{špatně } i} D_t(i)$, proto pokud chceme výraz minimalizovat přes h_t , musíme minimalizovat $\sum_{\text{špatně } i} D_t(i)$ (protože to má kladný koeficient $e^{a^*} - e^{-a^*}$).

d)

$$\begin{aligned} \text{Derivace podle } a &= e^{-a^*} * (-1) * \sum_{\text{všechny } i} D_t(i) + (e^{a^*} * 1 - e^{-a^*} * (-1)) * \sum_{\text{špatně } i} D_t(i) = \\ &= -e^{-a^*} * \sum_{\text{dobře } i} D_t(i) + e^{a^*} * \sum_{\text{špatně } i} D_t(i) = -e^{-a^*} * (1 - \varepsilon_t) + e^{a^*} * \varepsilon_t = 0 \end{aligned}$$

$$e^{-a^*} * (1 - \varepsilon_t) = e^{a^*} * \varepsilon_t$$

$$(1 - \varepsilon_t) / \varepsilon_t = e^{2a^*}$$

$$2a_t = \ln((1 - \varepsilon_t) / \varepsilon_t)$$

$$a_t = \frac{1}{2} * \ln((1 - \varepsilon_t) / \varepsilon_t)$$

e)

$$a_t = \frac{1}{2} * \ln((1 - \varepsilon_t) / \varepsilon_t) > 0 \text{ iff } \ln((1 - \varepsilon_t) / \varepsilon_t) > 0 \text{ iff } (1 - \varepsilon_t) / \varepsilon_t > 1 \text{ iff } \varepsilon_t < \frac{1}{2}$$

Problem 8.4

- 1) Inicializujeme váhy všech $n = 9$ bodů na $D_1(i) = 1/n = 1/9$ pro všechny body
- 2) nalezneme klasifikátor s nejmenší chybou $\varepsilon_1 = \sum_{\text{špatně } i} D_1(i)$
Tipuju, že to bude třeba $h_1 = \text{sign}(1 * x - 3.5)$, neboli oddělí body mezi 3 a 4, pak $\varepsilon_1 = 2/9$
- 3) Váha toho klasifikátoru a_1 je $\frac{1}{2} * \ln((1 - \varepsilon_1) / \varepsilon_1) = \frac{1}{2} * \ln(7/2) = 0,6264 = \ln(\sqrt{7/2})$
- 4) Přepočítáme váhy $D_2(i) = D_1(i) * e^{-a_1 h_1(x_i) y_i} / Z_1$, kde $Z_1 = \sum_{1 \leq j \leq 9} D_1(j) * e^{-a_1 h_1(x_j) y_j}$
Tj. pro špatně ohodnocené body 0 a 6:
 $D_1(j) * e^{-a_1 h_1(x_j) y_j} = 1/9 * e^{0,6264} = 1/9 * \sqrt{7/2} = 0,208$
A pro správně ohodnocené body:
 $D_1(j) * e^{-a_1 h_1(x_j) y_j} = 1/9 * e^{-0,6264} = 1/9 * \sqrt{2/7} = 0,06$
 $\Rightarrow Z_1 = 2 * 0,208 + 7 * 0,06 = 0,836$
A pak pro špatně ohodnocené body 0 a 6: $D_2(j) = 0,208 / 0,836 = 0,2488$
A pro správně ohodnocené body: $D_2(j) = 0,06 / 0,836 = 0,0718$
Lze si výpočet zjednodušit rovností $Z_t = 2 * \sqrt{(1 - \varepsilon_t) * \varepsilon_t}$
- 5) Pokračovalo by se znovu krokem 2)

Problem 9.1

- a) Backpropagation - 5 - Computationally efficient automatic differentiation for scalar-valued composite functions.
- b) Gradient - 3 - Is necessary to find a step direction for gradient descent.
- c) Chain rule - 4 - A rule to compute gradient of composite functions.
- d) Training loss minimization - 1 - A way to learn neural networks.
- e) SGD (stochastic gradient descent) - 2 - Method to optimize training loss.

Problem 9.2

- a) Gradient of f - 3 - Column vector of partial (or total) derivatives in case f is scalar-valued, i.e. $m = 1$.
- b) Derivative of f - 1 - A linear mapping approximating f locally around a point.
- c) Jacobian of f - 2 - Expression of the derivative in coordinates as a matrix.

Problem 9.3

Jakože po roce od MA2 si představovat, jak vlastně vypadají derivace ve vyšších dimenzích, je fakt záhul, fuj...

a)

$y(x) = Wx: \mathbb{R}^D \rightarrow \mathbb{R}^M$, $L(y): \mathbb{R}^M \rightarrow \mathbb{R}^1$, $L(y(x)): \mathbb{R}^D \rightarrow \mathbb{R}^1$, takže matice $\partial L / \partial x$ bude z $\mathbb{R}^{1 \times D}$

$$\partial L / \partial x_i = \sum_j \partial L / \partial y_j * \partial y_j / \partial x_i = \sum_j \partial L / \partial y_j * W_{j,i}$$

$\partial L / \partial x = \sum_j \partial L / \partial y_j * \partial y_j / \partial x = \sum_j \partial L / \partial y_j * W_{j,:} = \partial L / \partial y * W$, kde $\partial L / \partial y$ je z $\mathbb{R}^{1 \times M}$ (řádkový vektor) a W je z $\mathbb{R}^{M \times D}$

b)

Gradient je D^T . $D = \partial L / \partial y * W \rightarrow D^T = (\partial L / \partial y * W)^T = W^T * (\partial L / \partial y)^T = W^T * \text{Grad}_y L$

c) ?????

Bereme teď $L(W): \mathbb{R}^{M \times D} \rightarrow \mathbb{R}$

$$\partial L / \partial W_{i,j} = \partial L / \partial y_i * \partial y_i / \partial W_{i,j} = \partial L / \partial y_i * x_j$$

Protože je obecně derivace funkce $\mathbb{R}^a \rightarrow \mathbb{R}^b$ matice z $\mathbb{R}^{a \times b}$ (ne $\mathbb{R}^{b \times a}$), tak v matici $D = \partial L / \partial W$ bude na pozici $D_{i,j}$ hodnota $\partial L / \partial W_{i,j}$. Získáme:

$D = \partial L / \partial W = x * \partial L / \partial y$, kde x je z $\mathbb{R}^{D \times 1}$ a $\partial L / \partial y$ je z $\mathbb{R}^{1 \times M}$, tedy D je z $\mathbb{R}^{D \times M}$.

$\text{Grad}_W L = (\partial L / \partial W)^T = \text{Grad}_y L * x^T$ z $\mathbb{R}^{M \times D}$. Snad.

Tady je hezky vysvětlené: <http://cs231n.stanford.edu/handouts/linear-backprop.pdf>

Problem 10.1

a)

máme fixní $\{c_k\}_{k=1,\dots,K} \Rightarrow$ sumu minimalizujeme minimalizací jejích sčítanců $\Rightarrow$ pro všechny body x_i nalezneme nejbližší c_k : $r(i) = \min_{k=1,\dots,K} \|x_i - c_k\|^2$

b)

máme fixní dělení na clustery $\Rightarrow$ sumu vzdáleností od center ve všech clusterech minimalizujeme minimalizací vzdáleností od centra v jednotlivých clusterech $\Rightarrow$ v každém clusteru k uděláme: $c_k = \operatorname{argmin}_{c \in \mathbb{R}^D} \sum_{i: r(i)=k} \|x_i - c_k\|^2$, derivací získáme

$c_k = (\sum_{i: r(i)=k} x_i) / (\sum_{i: r(i)=k} 1)$, což je těžiště clusteru

Existuje-li centrum k takové, že $r(i) \neq k$ pro žádné i (neboli není mu přidělen žádný bod), pro jeho vrácení do koloběhu minimalizace ho reinitializujeme

c)

V K-means se střídají tyto dva kroky.

V prvním kroku si žádný bod nemůže pohoršit (zvětšit svou vzdálenost od nejbližšího centra vzhledem ke stavu před tímto krokem), tedy se nemůže zhoršit ani celková suma.

V druhém kroku se nemůže zhoršit součet vzdáleností v žádném clusteru, tedy se nemůže zhoršit ani celková suma (jednotlivé konkrétní body si pohoršit můžou, ale globálně ne).

Následující krok K-means tedy nemůže mít horší sumu, než předchozí. Pokud je lepší, pokračujeme, pokud je stejná, nic se nezměnilo (pokud je konzistentní rozhodování v případě rovnosti) a algoritmus končí v (lokálním) minimu $\Rightarrow$ nejde se zacyklit.

Celkových možných stavů je N^K , kde N i K jsou konečné, tedy algoritmus musí skončit v konečném počtu kroků.

Problem 10.2 ?????

Intuice říká, že čím menší je suma vzdáleností mezi dvojicemi bodů, tím menší je suma vzdáleností od těžiště (body jsou prostě "blíže" u sebe). Ale jak to matematicky ukázat na tom vzorci nevím.

Problem 10.3

a)

Hledáme pozice studní $c_k = (x_k, y_k)$, $k = 1, \dots, K$. Jako $d(x, y)$ použijeme euklidovskou vzdálenost (L2 normu bez kvadrátu):

$$\{c_k^*\} = \operatorname{argmin}_{\{c_k\} \in \mathbb{R}^{K \times 2}} \sum_{i=1, \dots, N} \min_k \|p_i - c_k\|$$

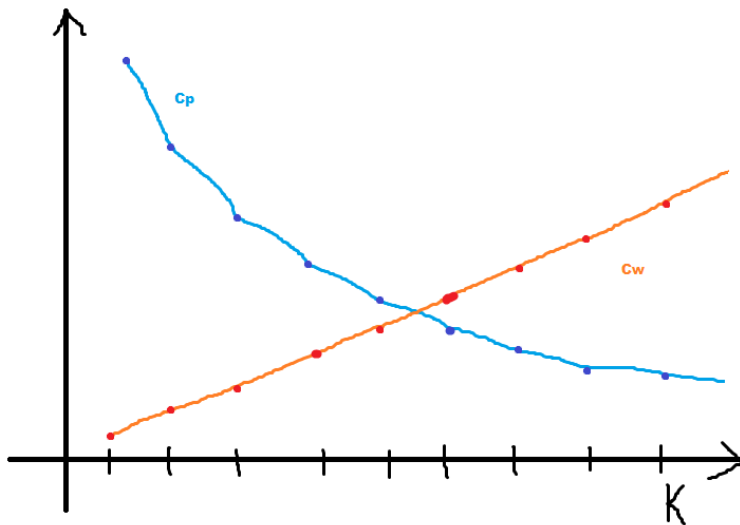
K-medians by používaly Manhattanskou L1 normu:

$$\{c_k^*\} = \operatorname{argmin}_{\{c_k\} \in \mathbb{R}^{K \times 2}} \sum_{i=1, \dots, N} \min_k \|p_i - c_k\|_1 = \operatorname{argmin}_{\{c_k\} \in \mathbb{R}^{K \times 2}} \sum_{i=1, \dots, N} \min_k (|p_{i1} - c_{k1}| + |p_{i2} - c_{k2}|)$$

b)

Nyní musíme brát v potaz kombinaci klesající ceny za potrubí a rostoucí ceny za studně.

Formálně: $K^* = \operatorname{argmin}_{K \in \{1, 2, \dots\}} (K^* c_w + \operatorname{argmin}_{\{c_1, \dots, c_K\}} \sum_{i=1 \dots N} \min_{j \in \{1, \dots, K\}} d(p_i, c_j) c_p)$



Vztah c_w vzhledem ke k je lineární. Vztah c_p není lineární, přijde mi, že platí, že každé další zlepšení hodnoty c_p s rostoucím k je menší než předchozí zlepšení (pokud by $i+1$. iterace našla bod X , který oproti i . iteraci udělá lepší změnu, než udělal bod Y nalezený v i . iteraci oproti $i-1$. iteraci, tak tento bod X mohl být nalezen již v i . iteraci místo bodu Y a udělat větší zlepšení, tedy neplatí, že Y je optimální řešení v i . iteraci, což je spor, mi přijde, ale případně mě vyvedte z omylu). Proto jakmile poprvé narazíme na iteraci m , ve které se součet $c_w + c_p$ zvětší, předchozí počet studní je ten optimální, tj. $k = m - 1$.

c)

Algoritmus pouze přijme jinou formu vzdálenosti - Manhattanskou vzdálenost (L1 normu):

$$\{c_k^*\} = \operatorname{argmin}_{\{c_k\} \in \mathbb{R}^{K \times 2}} \sum_{i=1, \dots, N} \min_k \|p_i - c_k\|_1 = \operatorname{argmin}_{\{c_k\} \in \mathbb{R}^{K \times 2}} \sum_{i=1, \dots, N} \min_k (|p_{i1} - c_{k1}| + |p_{i2} - c_{k2}|)$$

Problem 11.1

Máme $T = \{x_i\}_{i=1..N}$, $p(x|k) = N(\mu_k, I)$, $p(k) = 1/K$, $p(x, k) = p(x|k)p(k) = N(\mu_k, I)/K$, kde K je konstantní, proto ji můžeme v optimalizaci vynechat a brát $p(x, k) = N(\mu_k, I)$

a)

Nejdříve nalezneme pro každý bod x_i jeho odhadovanou třídu k_i :

$$k_i = \operatorname{argmax}_k p(x_i|k)p(k)/(\sum_{k'} p(x_i|k')p(k')) = \operatorname{argmax}_k p(x_i|k) = \operatorname{argmax}_k N(x_i | \mu_k, I)$$

Poté pro každou třídu zvlášť upravíme její μ_k :

$$\mu_k = \operatorname{argmax}_\mu (\operatorname{Prod}_{i: k_i = k} 1/\sqrt{(2\pi)^D} \exp(-\frac{1}{2} (x_i - \mu)^T (x_i - \mu)))$$

$$l_k(\mu) = \sum_{i: k_i = k} (\ln(1/(2\pi)^{D/2}) - \frac{1}{2} (x_i^T x_i - 2\mu^T x_i + \mu^T \mu))$$

$$\partial l_k(\mu)/\partial \mu = \sum_{i: k_i = k} (-x_i^T + \mu^T) = 0, \text{ odtud: } \mu_k = (\sum_{i: k_i = k} x_i^T)/(\sum_{i: k_i = k} 1)$$

b)

Nyní budeme mít pro každý bod váhy $\alpha_i(k) = p(k|x_i)$ pro všechny k .

$$\alpha_i(k) = p(x_i|k)p(k)/(\sum_{k'} p(x_i|k')p(k')) = p(x_i|k)/(\sum_{k'} p(x_i|k'))$$

Parametry μ_k můžeme upravovat nezávisle pro každou k :

$$l_k(\mu) = \sum_{i=1..N} (\alpha_i(k) \ln(\exp(-\frac{1}{2} (x_i - \mu)^T (x_i - \mu)))) = \sum_{i=1..N} \alpha_i(k) * (-\frac{1}{2} (x_i^T x_i - 2\mu^T x_i + \mu^T \mu))$$

$$\partial l_k(\mu)/\partial \mu = \sum_{i=1..N} (\alpha_i(k) * (-x_i^T + \mu^T)) = 0$$

$$\sum_{i=1..N} (\alpha_i(k) * x_i^T) = \sum_{i=1..N} (\alpha_i(k) * \mu^T) = \mu^T * \sum_{i=1..N} \alpha_i(k)$$

$$\mu^T = \sum_{i=1..N} (\alpha_i(k) * \mu^T) / \sum_{i=1..N} \alpha_i(k), \text{ což je vážený průměr}$$

Problem 12.1

a)

$$\begin{aligned}\text{Sum}_{j=1..M} (1/N * \text{Sum}_{i=1..N} (u_j^T x_i)^2) &= \text{Sum}_{j=1..M} (1/N * \text{Sum}_{i=1..N} u_j^T x_i u_j^T x_i) = \\ \text{Sum}_{j=1..M} (1/N * \text{Sum}_{i=1..N} u_j^T x_i x_i^T u_j) & \text{ (protože } u_j^T x_i \text{ je skalár) } = \\ \text{Sum}_{j=1..M} (1/N * u_j^T * (\text{Sum}_{i=1..N} x_i x_i^T) * u_j) &= \text{Sum}_{j=1..M} u_j^T S u_j\end{aligned}$$

b)

$$\begin{aligned}\text{Sum}_{i=1..N} (x_i - \text{Sum}_{j=1..M} u_j^T x_i u_j)^T (x_i - \text{Sum}_{j=1..M} u_j^T x_i u_j) &= \\ \text{Sum}_{i=1..N} (x_i^T x_i - 2 * x_i^T * (\text{Sum}_{j=1..M} u_j^T x_i u_j) + (\text{Sum}_{j=1..M} u_j^T x_i u_j)^T (\text{Sum}_{j=1..M} u_j^T x_i u_j)) &= \\ \text{Sum}_{i=1..N} (x_i^T x_i - 2 * (\text{Sum}_{j=1..M} x_i^T u_j^T x_i u_j) + (\text{Sum}_{j=1..M} \text{Sum}_{h=1..M} u_j^T x_i^T u_j u_h^T x_i u_h)) &= \dots \\ u_j^T x_i \text{ je skalár, můžeme to vytknout před vektory, máme pak } u_j^T x_i x_i^T u_j = u_j^T x_i u_j^T x_i = (u_j^T x_i)^2 & \\ x_i^T u_j \text{ a } u_h^T x_h \text{ jsou skaláry, vytkneme je před vektory: } x_i^T u_j u_h^T x_i u_j^T u_h, \text{ kde } u_j^T u_h = 1 \text{ pro } j = h, \text{ a} & \\ \text{jinak } u_j^T u_h = 0 \text{ (} u_i \text{ jsou ortonormální), takže získáme } \text{Sum}_{j=1..M} x_i^T u_j u_j^T x_i = \text{Sum}_{j=1..M} (u_j^T x_i)^2, \text{ tak:} & \\ \dots = \text{Sum}_{i=1..N} (x_i^T x_i - 2 * (\text{Sum}_{j=1..M} (u_j^T x_i)^2) + (\text{Sum}_{j=1..M} (u_j^T x_i)^2)) &= \text{Sum}_i (||x_i||^2 - (\text{Sum}_j (u_j^T x_i)^2))\end{aligned}$$

Problem 12.2

Z 12.1 b) víme, že chceme maximalizovat $\text{Sum}_{j=1..M} \text{Sum}_{i=1..N} (u_j^T x_i)^2$ a z 12.1 a) víme, že $\text{Sum}_{j=1..M} \text{Sum}_{i=1..N} (u_j^T x_i)^2 = N * \text{Sum}_{j=1..M} u_j^T S u_j$, kde N je konstanta, proto maximalizujeme $\text{Sum}_{j=1..M} u_j^T S u_j$, kde chceme ortonormální u_j , což je ekvivalentní té podmínce ve 12.2

a)

$$\begin{aligned}\text{Lagrange}(u) = L(u) &= \max_{\{u\}} \text{Sum}_{j=1..M} u_j^T S u_j - \lambda_1 (u_1^T u_1 - 1) - \dots - \lambda_M (u_M^T u_M - 1) \\ \partial L(u) / \partial u_j &= 2 u_j^T S - 2 \lambda_j u_j^T = 0, \text{ odkud } (S = (\text{Sum}_{i=1..N} x_i x_i^T) / N \text{ je symetrická):} \\ S u_j &= \lambda_j u_j\end{aligned}$$

b)

Rovnost $S u_j = \lambda_j u_j$ přímo popisuje vlastní čísla a vektory z definice a navíc vlastní vektory splňují ortogonalitu, respektive i ortonormalitu po normalizaci.

c)

Víme $S u_j = \lambda_j u_j$, takže lze maximalizaci přepsat jako:

$$\begin{aligned}\max_{\{u_j\}} \text{Sum}_{j=1..M} u_j^T S u_j &= \max_{\{u_j\}} \text{Sum}_{j=1..M} u_j^T \lambda_j u_j = \max_{\{u_j\}} \text{Sum}_{j=1..M} \lambda_j u_j^T u_j = \max_{\{u_j\}} \text{Sum}_{j=1..M} \lambda_j, \\ \text{kde tedy } \lambda_j &\text{ je vlastní číslo vlastního vektoru } u_j. \text{ Abychom to maximalizovali, bereme tedy } M \\ \text{vlastních vektorů patřících k } M &\text{ největším vlastním číslům.}\end{aligned}$$

Problem 12.3

Proč vlastně minimalizujeme tohle? Původní úloha byla maximalizovat:

$(\mu_1 - \mu_2)^2 / (s_1 + s_2)$, kde μ_k je střední hodnota projektovaných dat třídy k a s_k je rozptyl projektovaných dat třídy k .

$\mu_k = 1/N_k * \text{Sum}_{x_i \text{ třídy } k} v^T x_i = v^T (1/N_k * \text{Sum}_{x_i \text{ třídy } k} x_i) = v^T x_k$, kde x_k je těžiště bodů třídy k
 $(\mu_1 - \mu_2)^2 = (v^T x_1 - v^T x_2)^2 = (v^T (x_1 - x_2))^2 = v^T (x_1 - x_2) v^T (x_1 - x_2)$, kde $v^T (x_1 - x_2)$ je skalár, lze tedy transponovat, proto:

$$\begin{aligned}(\mu_1 - \mu_2)^2 &= v^T (x_1 - x_2) (x_1 - x_2)^T v = v^T S_b v \\ s_k &= \text{Sum}_{x_i \text{ třídy } k} (v^T x_i - v^T x_k)^2 = \text{Sum}_{x_i \text{ třídy } k} (v^T (x_i - x_k))^2 = \text{Sum}_{x_i \text{ třídy } k} v^T (x_i - x_k) v^T (x_i - x_k), \text{ kde} \\ v^T (x_i - x_k) &\text{ je skalár, lze tedy transponovat, proto:} \\ s_k &= \text{Sum}_{x_i \text{ třídy } k} v^T (x_i - x_k) (x_i - x_k)^T v = v^T (\text{Sum}_{x_i \text{ třídy } k} (x_i - x_k) (x_i - x_k)^T) v = v^T S_k v\end{aligned}$$

Úloha se přepíše na:

$$\begin{aligned}\max (\mu_1 - \mu_2)^2 / (s_1 + s_2) &= \max_v v^T S_b v / (v^T S_1 v + v^T S_2 v) = \max_v v^T S_b v / v^T (S_1 + S_2) v = \\ \max_v v^T S_b v / v^T S_w v\end{aligned}$$

Pokud je S_w regulární, existuje S_w^{-1} , definujeme $z = S_w^{-1/2} v$, a získáme:

$$\max z^T S_w^{-1/2} S_b S_w^{-1/2} z / z^T z, \text{ kde } z^T z = v^T S_w v = 1, \text{ takže máme:}$$

$$\max z^T S_w^{-1/2} S_b S_w^{-1/2} z \quad \text{za podmínky } z^T z = 1$$

$$L(z) = z^T S_w^{-1/2} S_b S_w^{-1/2} z - \lambda(z^T z - 1)$$

$$\partial L(z)/\partial z = 2 * z^T S_w^{-1/2} S_b S_w^{-1/2} - 2 * \lambda z^T = 0$$

$$S_w^{-1/2} S_b S_w^{-1/2} z = \lambda z$$

$$S_w^{-1/2} S_b S_w^{-1/2} S_w^{1/2} v = \lambda S_w^{1/2} v$$

$$S_w^{-1} S_b v = \lambda v$$

$$S_b v = \lambda S_w v$$

Problém hledání vlastního čísla byl předchozí krok $S_w^{-1} S_b v = \lambda v$, kdy hledáme vlastní číslo a vektor matice $S_w^{-1} S_b$.

$\max_v v^T S_b v / v^T S_w v = \max_v \lambda v^T S_w v / v^T S_w v = \max_v \lambda$, kde λ je vlastní číslo vlastního vektoru v matice $S_w^{-1} S_b$, neboli volíme vlastní vektor matice $S_w^{-1} S_b$ příslušící jejímu největšímu vlastnímu číslu.

Problem 12.4

Střední hodnota:

$$\mu_k = 1/N_k * \sum_{x_i \text{ třídy } k} v^T x_i = v^T (1/N_k * \sum_{x_i \text{ třídy } k} x_i) = v^T x_k$$

Rozptyl:

$$s_k = \sum_{x_i \text{ třídy } k} (v^T x_i - v^T x_k)^2 = \sum_{x_i \text{ třídy } k} (v^T (x_i - x_k))^2 = \sum_{x_i \text{ třídy } k} v^T (x_i - x_k) v^T (x_i - x_k), \text{ kde}$$

$v^T (x_i - x_k)$ je skalár, lze tedy transponovat, proto:

$$s_k = \sum_{x_i \text{ třídy } k} v^T (x_i - x_k) (x_i - x_k)^T v = v^T (\sum_{x_i \text{ třídy } k} (x_i - x_k) (x_i - x_k)^T) v = v^T \Sigma_k v$$

Semestrální test 2 (2013)

https://cw.fel.cvut.cz/b201/media/courses/be5b33rpz/labs/rpzpractice_test2.pdf

Problem 1: ($h = \lambda$)

Poisson(h): $p(x) = h^x e^{-h} / x!$

$L(h) = p(X|h) = \prod_{i=1..n} (h^{x_i} e^{-h} / x_i!)$

$l(h) = \sum_{i=1..n} (x_i \ln(h) - \ln(x_i!)) - nh$

$l'(h) = \sum_{i=1..n} (x_i/h - 1) = 0$, odtud $h = \sum_{i=1..n} x_i / n = \text{ar. průměr}$

Problem 2:

Exp(h): $p(x) = h^x e^{-hx}$

Obece h_{MLE} :

$L(h) = p(X|h) = \prod_{i=1..n} (h^{x_i} e^{-hx_i})$

$l(h) = \sum_{i=1..n} (\ln(h) - h x_i)$

$l'(h) = \sum_{i=1..n} (1/h - x_i)$, odtud $h = n / (\sum_{i=1..n} x_i) = (\text{ar. průměr})^{-1}$

Konkrétně je to $h_{MLE} = 5 / 13$

Obece dele než 5 let: (asi, nejsem si jistý, ale vychází to smysluplně)

$p(x > 5) = \int_{5..inf} p(x) dx = \int_{5..inf} h^x e^{-hx} dx = [h^x (-1/h) e^{-hx}]_{5..inf} =$
 $= [-1^x e^{inf} + e^{-5h}] = e^{-5h}$

Konkrétně to je $p(x > 5) = e^{-5 \cdot 5/13} = e^{-25/13} = 0,146...$

Problem 3:

Je jedno, ze jsou ty sekvence nějak deleny, prostě:

$T = (\text{HHTHHTTTHHHHTHHHTH})$, $n = 17$, $n_H = 12$ a $n_T = 5$

Intuice: $p = p(H) = n_H/n = 12/17$, $1 - p = p(T) = 5/17$

$p_{MLE} = \text{argmax } p(X|p)$

$L(p) = p(X|p) = \prod_{k \text{ from } T} p(k) = p(H)^{n_H} p(T)^{n_T}$

$l(p) = n_H \ln(p) + n_T \ln(1 - p)$

$l'(p) = n_H/p - n_T/(1 - p) = 0$, odtud: $n_H/p = n_T/(1 - p)$

$n_H(1 - p) = n_T p$

$n_H - n_H p = n_T p$

$n_H = n_T p + n_H p$

$p_{MLE} = n_H / (n_H + n_T) = n_H / n = 12/17$

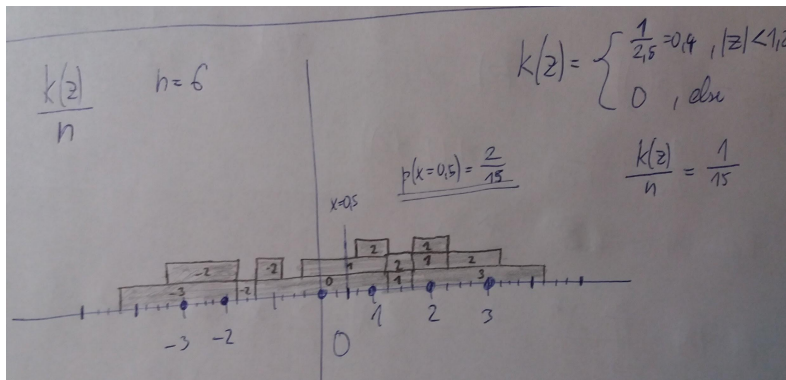
Problem 4:

Jadrova funkce:

$k(z) = 1/(2 \cdot 1,25) = 2/5$ iff $|z| \leq 1.25$

$k(z) = 0$ iff $|z| > 1.25$

Do každého bodu vložíme $k(z)/n$, získáme $p(x = 0.5) = 2/15$



Problem 5: ?????

Ani nevím, jestli tu otázku poradne chapu, je to “vyhodnotili jste kvalitu MLE odhadu tím, že jste vykreslili jejich rozptyl jako funkci velikosti trénovací množiny”? Nejak nevím co odpovedet.

Problem 6: ?????

Zas mi ta otázka přijde nejaka divna, tak zalezi co jsme odhadovali, ne? Jinak odhadnout, ze neco je z normalniho rozdeleni mi přijde jako fajn volba, hodne veci ve svete odpovida normalnimu rozdeleni. Co tu ale rict konkretně fakt nevím.

Poznámka ze cvika:

z centralni limitni vety víme, že pro nekonečne mnoho vzorku se vlastnosti spojitých rozděléní blíží normalnimu.

Semestrální test 3 (ukázka)

https://cw.fel.cvut.cz/b191/courses/be5b33rpz/labs/exercises/test3practice?fbclid=IwAR0k0FrWOSWAuiCXaAZ8IS_O2bQM3goDj7Jk4pLHfH-SnpurQG5z_RdMJU

Task 1

Máme trénovací množinu T s body $x' \in \mathbb{R}^{D+1}$, které vznikly z bodů $x \in \mathbb{R}^D$ a jejich tříd $y = \{-1, 1\}$, jako $x' = (1, x)^T y$

- 1) Inicializujeme $t = 0$ (číslo iterace), $w_t = 0 \in \mathbb{R}^{D+1}$
- 2) Nalezneme špatně ohodnocené $x \in T$, tj. takové x , že $w_t^T x \leq 0$. Pokud takové x neexistuje, končíme a $w = w_t$ vektor popisující oddělovací nadrovinu.
- 3) Updatujeme $w_{t+1} = w_t + x$ a jdeme na krok 2).

Se získaným vektorem w poté klasifikujeme takto: $q(x) = \text{sign}(w^T x)$, kde $x \in \mathbb{R}^{D+1}$ je testovaný vektor upravený stejně jako trénovací.

Task 2

Perceptron nalezne jakoukoli oddělovací nadrovinu a poté končí. SVM nalezne nejlepší oddělovací nadrovinu ve smyslu velikosti margin - je to oddělovací nadrovinu nejvzdálenější od nejbližších bodů. To znamená, že SVM lépe generalizuje (má nižší VC dimenzi).

Task 3

Trénovací data jsou (x, y) , C je ztráta za špatně klasifikované x (poměrná ke vzdálenosti)

Primární úloha:

$$w^*, b^* = \arg\min_{w,b} \frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad \text{za podmínky } y_i(w^T x_i + b) \geq 1 - \xi_i, \xi_i \geq 0$$

Duální úloha:

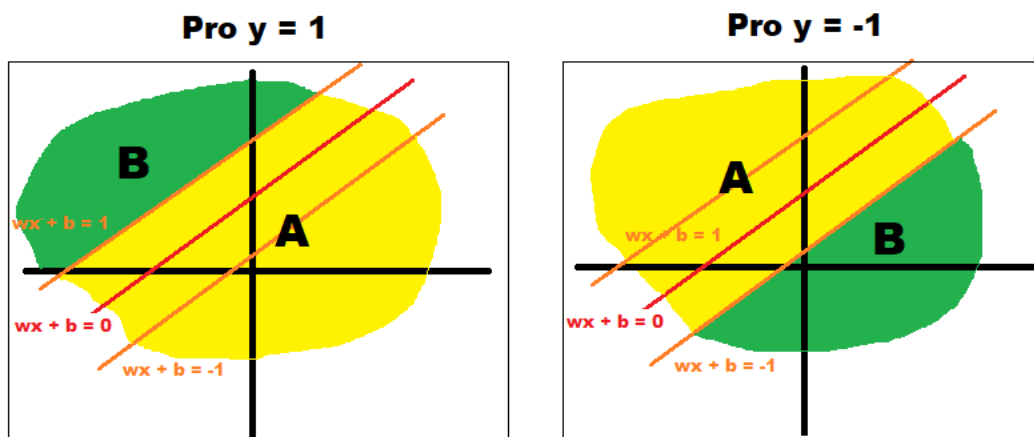
$$a = \arg\max_a \sum_i a_i - \frac{1}{2} \sum_i \sum_j a_i a_j y_i y_j x_i^T x_j \quad \text{za podmínky } 0 \leq a_i \leq C, \sum_i a_i y_i = 0$$

$$w = \sum_i a_i y_i x_i$$

$$b = y_s - w^T x_s \quad \text{pro nějaký support vector } s$$

Klasifikujeme jako $q(y) = \text{sign}(w^T y + b)$

Task 4



Task 5 - Nejsem si stopro jistý, opravujte.

Očekávaná ztráta $R(q)$ je průměrná ztráta na testovacích datech. Potřebujeme k ní znát sdruženou pravděpodobnost tříd a pozorování. $R(q) = \sum_{k,x} p(x, k) W(q(x), k)$

Empirická ztráta je ztráta na trénovacích datech. Lze vypočítat bez znalosti sdružené pravděpodobnosti. $R_{\text{emp}}(q) = 1/L \sum_{i=1}^L W(q(x_i), k_i)$

Upřednostníme očekávanou ztrátu kdykoli známe $p(x, k)$, nebo lze dobře odhadnout.

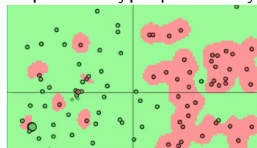
Empirická ztráta má širší využití.

Semestrální test 3

Úloha 1. (2b.) Máme trénovací sadu s měřeními $x_i \in \mathbb{R}^2$ a labely $y_i \in \{1, 2\}$, $T = \{((3, 1); 2), ((2, 2); 2), ((2, 1); 1), ((1, 3); 1)\}$. Proveďte 3 iterace učení Perceptronu (vypisujte důležité hodnoty).

Úloha 2. (1b.) Trénujete soft-margin SVM s Gauss kernelom $K(x, y) = \exp(-\|x - y\|^2/2\sigma^2)$. Klasifikátor na tréningových dátach je na obrázku, na validačných dátach ale nefunguje dobre. Jak upravíte hyperparametry C a σ ? Proč?

1. zmenšit C , zmenšit σ
2. zmenšit C , zvětšit σ
3. zvětšit C , zmenšit σ
4. zvětšit C , zvětšit σ



Úloha 3. (1b.) Jaká vlastnost platí pro vztah mezi optimálními hodnotami kriteriální funkce primální a duální úlohy SVM?

Úloha 4. (1b.) Mějme funkci $K_\alpha(x, y) = \alpha(x + y)$ pro $x, y \in \mathcal{R}^3$ a $\alpha > 0$. Je tato funkce validní jádrová funkce (kernel)? Pokud ne, dokažte. Pokud ano, najděte odpovídající feature mapu $\phi(x)$.

① - danine ei body na $y_i \in \{-1, 1\}$, β : $2 \Rightarrow -1$
 - njeanthen T ~~kl. kl.~~: $\bar{x}_i = y_i \cdot (x_i, 1) \in \mathbb{R}^3$
 - wiskun $T = \{(-3, -1, -1), (-2, -2, -1), (2, 1, 1), (1, 3, 1)\}^T$
 o) inisializazjone $w_0 = (0, 0, 0)$
 1) β pru odothoweni $x_1 = (-3, -1, -1) \Rightarrow w_1 = w_0 + x_1 = (-3, -1, -1)$
 2) ---||--- $x_2 = (2, 1, 1) \Rightarrow w_2 = w_1 + x_2 = (-1, 0, 0)$
 3) ---||--- $x_3 = (2, 1, 1) \Rightarrow w_3 = w_2 + x_3 = (1, 1, 1)$
 (β pru odothoweni x iff $w_i^T x \leq 0$)
 - β body june w ilowei hledeal β pru odothoweni x od roatitku trawowei mawiny)
 - na lowei bjeowen rishali $w_i = (a, b, c)$ a odditijic madowow (grawow) w grawowow
 dalek w $\mathbb{R}^2$ by lje $a \cdot x_1 + b \cdot x_2 + c = 0$

② - zvolit bych mohl 2: zmenšime C tak, abychom provolili ∇ nějaké bodové množiny bodů, a zvolíme rovněž, at' přines bodu do klasifikace nové bodě charakterizující

- klasifikace je ^{tedy} grafický

③ - dualitate mai înseamnă, în optimizării liniare, că avem soluții de ambele probleme.

④ - $K_\alpha(x, y) = \alpha(x+y)$ nemie je kernel funkcia, pretože reciproci stolar, ale vedie, kernel funkcia obratú je $K: \mathbb{R}^D \times \mathbb{R}^D \rightarrow \mathbb{R}$.

(i) Bilo $\alpha \in \mathbb{R}^3$, $\alpha_i > 0$, tak $\alpha^T(x+y)$ nema pozitivnu semidefinitivost.
 (ovo je lakše pokazivati na konkretnim funkcijama) ↑
može biti ~~određeni~~ skalar $\in \mathbb{R}$

Semestrální test 3

1. Write down the dual task criterion function optimised by soft-margin SVM and define all symbols (only the numeric solution will be awarded by a point). How do you recognize the support vectors? (1 point)
2. Define a kernel function $K(\mathbf{x}, \mathbf{y})$. Discuss one advantage of using the kernel Support Vector Machine compared to the linear SVM. (1 point)
3. Write down three advantages of the soft-margin SVM in comparison with the Perceptron algorithm. (1 point)
4. You are given a linear classifier $\mathbf{w} \in \mathbb{R}^{D+1}$ and a training dataset $\mathcal{T} = \{(\mathbf{x}_1; k_1), (\mathbf{x}_2; k_2), \dots, (\mathbf{x}_L; k_L)\}$, where $\mathbf{x}_i \in \mathbb{R}^D$, $k_i \in \{-1, 1\}$ with $i = 1, 2, \dots, L$. Formally define the linear separability. Only the numeric solution will be awarded by a point. (1 point)
5. You are given a training dataset $\mathcal{T} = \{(-1, 0.5; 1), (3, 3; 1), (2, -2; 1), (-2, -2; -1), (-1, -1; -1), (-1, -2.5; -1)\}$, where $\mathbf{x}_i \in \mathbb{R}^2$, $k_i \in \{-1, 1\}$ with $i = 1, 2, \dots, 6$. Does the perceptron algorithm finish in a finite number of steps on this dataset? If so, why (give a numeric solution)? (1 point)

1)

$\alpha^* = \operatorname{argmax}_{\alpha} \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j \mathbf{x}_i^T \mathbf{x}_j$ za podmínky $\sum_i \alpha_i y_i = 0$, $0 \leq \alpha_i \leq C$
 α jsou Lagrangeovy multiplikátory, použijí se jako váhy trénovacích vektorů ve finálním klasifikátoru: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$, $b = y_{sv} - \mathbf{w}^T \mathbf{x}_{sv}$ (pro nějaký support vektor $\mathbf{x}_{sv}$)

$(\mathbf{x}_i, y_i)$ jsou trénovací body $\mathbf{x}_i \in \mathbb{R}^D$ a jejich třídy $y_i = \pm 1$

C je ztráta za špatně klasifikovaný bod (úměrná vzdálenosti od správně klasifikovaného poloprostoru)

Support vektory jsou takové $(\mathbf{x}_i, y_i)$ z trénovací množiny, které mají nenulové váhy α

2)

$K(\mathbf{x}, \mathbf{y}) = \phi(\mathbf{x})^T \phi(\mathbf{y})$, kde ϕ je mapování do vyšší dimenze $\phi: \mathbb{R}^D \rightarrow \mathbb{R}^M$, které umožňuje použít lineární klasifikaci u lineárně neseparovatelných dat.

Výhoda: Lze tedy použít lineární klasifikátor pro oddělení lineárně neseparovatelných dat i bez explicitní znalosti ϕ . To umožňuje použít i potenciálně nekonečně dimenzionální mapování.

3)

- a) umí oddělit lineárně neseparovatelná data
- b) najde optimální řešení z hlediska velikosti marginu (v poměru ke ztrátě za špatnou klasifikaci C), narozdíl od nějakého řešení u Perceptronu
- c) při učení využívá pouze skalární součiny trénovacích dat, je tedy nezávislé na dimenzi

4)

Množina T je lineárně separabilní, pokud existuje $\mathbf{w} \in \mathbb{R}^{D+1}$ takové, že pro všechny $(\mathbf{x}, k)$ z T platí: $\operatorname{sign}(\mathbf{w}^T(\mathbf{x}, 1)) = k$

5) ?????

Jediný, co mě napadá, jak to zjistit, je si to nakreslit, nevím, jaký chtějí numerické řešení.

Novikoff theorem nám nepomůže, protože ten pouze říká "když máme vzdálenost γ a maximální velikost vektoru D , tak pokud jsou data separovatelná, tak skončí do iterace $t^* = D^2/\gamma^2$ ". Takže i kdybychom γ a D našli, tak to nic neříká o obecném zastavení algoritmu, když nevíme, jestli je to separovatelné.

4. RPZ Test 16.12.2020

Task 1 (AdaBoost) Consider the training set $\{(x_i, y_i) \mid i = 1 \dots n\}$ where $y_i \in \{-1, 1\}$. The AdaBoost strong classifier at step t is given by $H_t(x) = \text{sign}(f_t(x))$, where $f_t(x) = \sum_{k=1}^t \alpha_k h_k(x)$, $\alpha > 0$ and $h_k(x) \in \{-1, 1\}$.

a) Show that the number of training points misclassified by H_t is upper bounded by $\sum_{i=1}^n e^{-y_i f_t(x_i)}$. Properties of the exponent function e^z need not be proven.

b) To each data point i associate a weight $W_t(i)$ such that the upper bound in a) can be expressed as: $\sum_i W_t(i) e^{-\alpha_t y_i h_t(x_i)}$. It is not required that weights would sum to 1.

c) Show that if a training point (x_i, y_i) is misclassified by the weak classifier h_t , then its weight in the next iteration $t + 1$ will be increased and if it is correctly classified, then decreased.

(2 points)

Task 2 (K-means / EM) Let $\{x_i \mid i = 1 \dots n\}$ be observed points from a joint model $p(x, k)$ where $k \in \{1 \dots K\}$ is a hidden state (not observed). The model $p(x, k)$ is defined as $p(x, k) = p(x|k)p(k)$, where $p(x|k)$ is a Normal density $\mathcal{N}(\mu_k, I)$ with mean parameter μ_k for each k and $p(k)$ is uniform. Assume that initial estimates of μ_k are given.

a) Obtain K-means algorithm as follows. Given current means μ , find the most likely k_i , the assignment of mixture components to points maximizing the posterior distribution $p(k_i|x_i)$. Given current assignment k_i , find the maximum likelihood estimate of μ , i.e. maximizing the joint likelihood of x and k .

b) Obtain EM algorithm as follows. Given current means μ , find the soft assignment $\alpha_i(k) = p(k_i|x_i)$. Given weights $\alpha_i(k)$, find the parameters μ maximizing the “weighted” log-likelihood:

$$\sum_i \sum_k \alpha_i(k) \log p(x_i, k_i).$$

(3 points)

Task 1

a)

Špatně klasifikované body jsou ty, u kterých $H_t(x_i) \neq y_i$, neboli $f_t(x_i) \cdot y_i < 0$.

Platí $[H_t(x_i) \neq y_i] \leq e^{-y_i f_t(x_i)}$ pro všechny $(x_i, y_i) \in T$ (viz **Problem 8.1 b**), neboli platí horní mez i pro sumu:

$$\sum_{i=1, \dots, N} [H_t(x_i) \neq y_i] \leq \sum_{i=1, \dots, N} e^{-y_i f_t(x_i)}$$

b) ?????

$$W_t(i) e^{-\alpha_t h_t(x_i) y_i} = e^{-y_i f_t(x_i)} = e^{-y_i \sum_{j=1}^{t-1} \alpha_j h_j(x_i) - \alpha_t h_t(x_i) y_i} = e^{-y_i \sum_{j=1}^{t-1} \alpha_j h_j(x_i)} e^{-\alpha_t h_t(x_i) y_i}, \text{ odtud:}$$

$$W_t(i) = e^{-y_i \sum_{j=1}^{t-1} \alpha_j h_j(x_i)}$$

c)

Update váhy bodu (x_i, y_i) je definovaný $W_{t+1}(i) = W_t(i) \cdot e^{-\alpha_t h_t(x_i) y_i}$, kde $\alpha_t > 0$ je váha slabého klasifikátoru $h_t(x_i)$, $\alpha_t = \frac{1}{2} \ln((1 - \text{err}_t)/\text{err}_t)$ pro $\text{err}_t < \frac{1}{2}$.

U špatně klasifikovaný bodů je $h_t(x_i) \cdot y_i = -1$, neboli $W_{t+1}(i) = W_t(i) \cdot e^{\alpha_t}$, $e^{\alpha_t} > 1$.

U správně klasifikovaný bodů je $h_t(x_i) \cdot y_i = 1$, neboli $W_{t+1}(i) = W_t(i) \cdot e^{-\alpha_t}$, $e^{-\alpha_t} < 1$.

Task 2

Viz **Problem 11.1**

4. RPZ Test 18.12.2020

Úloha 1. (1b.) Jaká je časová složitost (počet float operací sčítání/násobení) základní verze algoritmu K-means pro K clusterů, T iterací a data $X = \{x_j \in \mathbb{R}^d \mid j = 1, \dots, N\}$?

Úloha 2. (1b.) Je u algoritmu K-means zaručená konvergence ke globálnímu optimu? Pokud ano, dokažte tuto vlastnost. Pokud ne, nakreslete příklad lokálního (ale ne globálního) optima.

Úloha 3. (1b.) Odvození klasifikátoru Adaboost je založeno na vztahu s empirickou chybou. Na jakém? (Odpovězte vzorcem nebo jednou větou.)

Úloha 4. (2b.) Uveďte vzorec z algoritmu Adaboost pro převážení trénovacích dat, tj. výpočet $D_{t+1}(i)$, a ukažte, jak zvyšuje váhu chybně klasifikovaných příkladů a snižuje váhu správně klasifikovaných příkladů.

① ~ *body iterací*: • nalezení nejbližšího clusteru c pro každý bod x :
 $O(K \cdot N \cdot d)$ (hladíme pro každý cluster) $\|x_i - c_k\|^2 = (x_{i1} - c_{k1})^2 + \dots + (x_{id} - c_{kd})^2$
 • upřesnění klíče clusterů:
 $O(N \cdot d)$ (pro každý cluster a pro každý jeho příslušný bod)
 $\Rightarrow$ celkově $O(T \cdot K \cdot N \cdot d)$

② *nový, konverguje k novému lokálnímu optimu*:
 $\Rightarrow$ pro $N=4$ body (A, B, C, D) a $K=2$
 clustery lze nalézt takto lokálně minimum
 c_1, c_2 (globální by byly c'_1, c'_2)

③ $\epsilon_{\text{emp}} = \frac{1}{N} \sum_{i: H_t(x_i) \neq y_i} 1 \leq \prod_{t=1}^T Z_t$, kde $Z_t = 2\sqrt{\epsilon_t(1-\epsilon_t)}$, kde ϵ_t je
 nejmenší chyba t -tého slabého klasifikátoru
 (nejmenší váhová kombinace dat $D_t(i)$)

④ $D_{t+1}(i) = \frac{D_t(i) \cdot e^{-\alpha_t \cdot h_t(x_i) \cdot y_i}}{Z_t}$, kde $Z_t = 2\sqrt{\epsilon_t(1-\epsilon_t)}$ je normalizační faktor,
 $\alpha_t = \frac{1}{2} \ln\left(\frac{1-\epsilon_t}{\epsilon_t}\right) > 0$ je váha t -tého slabého klasifikátoru a chyba $\epsilon_t < \frac{1}{2}$
 • pokud $h_t(x_i) = y_i$, pak $h_t(x_i) \cdot y_i = 1$, $e^{-\alpha_t \cdot 1} < 1 \Rightarrow$ sníží váhu správně klasifikovaných
 • pokud $h_t(x_i) \neq y_i$, pak $h_t(x_i) \cdot y_i = -1$, $e^{-\alpha_t \cdot (-1)} = e^{\alpha_t} > 1 \Rightarrow$ zvýší váhu špatně klasifikovaných
 $h_t(x_i) \in \{-1, 1\}$, $y_i \in \{-1, 1\}$

4. RPZ Test 18.12.2020

1. The criterion of AdaBoost algorithm has an interesting connection with the empirical error. What is their relation? One keyword in one short sentence, please. (1 point)
2. In iteration t of the Adaboost algorithm, we select the weak classifier with the lowest value of the Adaboost criterion $Z_t = \sum_i D_t(i) e^{-\alpha_t y_i h_t(x_i)}$. We then need to assign the weight α_t to the weak classifier. How do we find that weight? Give formal solution. (1 point)
3. Fill in the missing words. If the series of values of clustering criterion is monotonically _____ and convergent, and there is _____ number of partitions, and no assignment is visited twice, the K-means algorithm reaches _____ minimum. In the K-means algorithm, the computational time is dominated by the complexity of _____ step. (1 point)
4. Formally define the **generalized** K-means algorithm. Briefly explain all symbols. (1 point)
5. Is the K-Means algorithm guaranteed to find a global minimum? If so, explain why. If not, give an example of some local optimum setting. (1 point)

1)

Keyword: horní odhad

(přesněji: $1/N * \sum_{i: H(x_i) \neq y_i} 1 \leq \prod_{t=1..T} Z_t$, kde Z_t jsou normalizační konstanty rozdělení vah bodů $D_t(i)$, které minimalizujeme optimální volbou vah slabých klasifikátorů α_t)

2)

Derivací Z_t podle α_t :

$$Z_t = \sum_{i: H(x_i) \neq y_i} D_t(i) * e^{\alpha_t} + \sum_{i: H(x_i) = y_i} D_t(i) * e^{-\alpha_t} = \epsilon_t * e^{\alpha_t} + (1 - \epsilon_t) * e^{-\alpha_t}$$

$$\partial Z_t / \partial \alpha_t = \epsilon_t * e^{\alpha_t} - (1 - \epsilon_t) * e^{-\alpha_t} = 0$$

$$\epsilon_t * e^{\alpha_t} = (1 - \epsilon_t) * e^{-\alpha_t} \dots \alpha_t = \frac{1}{2} * \ln((1 - \epsilon_t) / \epsilon_t)$$

3)

decreasing, finite, local, "data-to-cluster reassignment"

4)

Obecný K-means:

Máme N bodů x_i a chceme je rozdělit do K shluků.

1) init centra ($c_1, \dots, c_K$) (např. náhodně nebo podle K-Means++)

2) nalezneme každému bodu nejbližší centrum, tj. vytvoříme shluky:

$T_k = \{x_i \mid d(x_i, c_k) \leq d(x_i, c_j) \text{ pro všechna } j\}$ (v případě rovnosti máme pevnou strategii rozhodování)

3) updatujeme centra:

$c_k = \operatorname{argmin}_c \sum_{x_i \in T_k} d(x_i, c)$, pokud T_k není prázdná, jinak reinit c_k

4) Pokud se T_k nezměnily, končíme, jinak pokračujeme krokem 2)

c_k, x_i jsou z R^D , $d: R^D \times R^D \rightarrow R$ je funkce vzdálenosti

5)

Nemusí, viz ten příklad s obdélníkem