



Národní konference s mezinárodní účastí
INŽENÝRSKÁ MECHANIKA 2002

13. – 16. 5. 2002, Svratka, Česká republika

Q-LEARNING: CONTROL OF ASYNCHRONOUS ELECTRIC MOTOR

Tomáš Marada¹

***Abstrakt:** Tento článek obsahuje numerické experimenty, popisující nastavení parametrů Q-učení pro řízení asynchronního elektromotoru. Cílem je stanovit velikost a strukturu stavového prostoru, vhodně zvolit optimální akce pro řízení a dále určit odměnu pro posilovací funkci. V závěru je provedeno porovnání s klasickým PID regulátorem.*

Klíčová slova

Asynchronní elektromotor, Q-učení, řízení

1. Úvod

Jednou z možností jak zlepšit kvalitu řízení asynchronního elektromotoru, je využití metod s netradičním řízením. Například se může jednat o implementaci metod umělé inteligence, speciálně těch metod, které používají strojové učení v reálném čase.

Metody opakovaně posilovaného učení (Reinforcement Learning) pracují na principu naučení agentova chování v dynamickém prostředí metodou „pokus-omyl“. To se totiž jeví jako přirozený způsob učení nejen pro zvířata, ale i pro člověka. Agent stejně jako živý organismus vnímá prostředí, ve kterém se právě nachází, provádí v něm různé akce a na základě jejich výsledků dostává číselné ohodnocení (Reinforcement) od vhodné posilovací funkce. Tento způsob učení má tu výhodu, že není nutné dopředu znát funkci, kterou je potřeba se naučit, protože nalezení vhodné funkce je neočekávaně těžkým úkolem, mnohdy neřešitelným.

Q-učení je jednou z nejpopulárnějších a nejpoužívanějších metod patřících do oblasti opakovaně posilovaného učení. Naším cílem je optimalizovat vztah mezi jakýmsi řídicím členem a dynamickým prostředím. V našem případě je agentem QL-regulátor a prostředím je asynchronní elektromotor chápaný jako 3-D svět se spojitými souřadnicemi. Úhel natočení elektromotoru, úhlová rychlost a úhlové zrychlení představuje jednu z možných konstrukcí stavů prostředí. Úkolem QL-regulátoru je postupně se naučit, prostřednictvím simulačního modelu asynchronního elektromotoru, minimalizovat hodnotu integrálního kritéria kvality regulace.

¹ Tomáš Marada, Ing., VUT v Brně, FSI ÚMT, Technická 2, 616 69 Brno, tomas.marada@seznam.cz

V algoritmu Q-učení se vyskytují následující pojmy:

Agent – v našem případě je agentem QL-regulátor. Úkolem agenta je, aby svým chováním (generováním akcí) dostal aktuální otáčky asynchronního elektromotoru do požadovaného intervalu otáček a v tomto intervalu je udržel. Agent tak představuje řídicí člen.

Model – je v našem případě výše zmíněný asynchronní elektromotor.

Akce – pomocí sady akcí agent řídí model a tím mění jeho aktuální stav v prostředí.

Stavový prostor – popisuje stavy prostředí ve kterých se může model nacházet. Na vhodnou volbu stavového prostoru závisí nejen kvalita naučení, ale i rychlost naučení. Pro stavový prostor je velmi důležitá jeho velikost, ale neméně významná je i jeho organizace. Z hlediska organizace stavového prostoru je významné, které veličiny jsou pro naučení podstatné. V naší úloze jsou podstatné tyto veličiny: Řídicí frekvence a aktuální otáčky. Samotná volba toho, co je dobré znát, však nestačí. Jelikož se jedná o spojitě veličiny, je třeba je vhodně rozdělit na určité intervaly. Čím více jich bude, tím přesněji bude systém fungovat, ale učení bude trvat déle. Stavový prostor lze rozdělit několika způsoby. Můžeme ho rozdělit na stejně velké intervaly, nebo na intervaly mající různou velikost. První jmenovaná možnost se nazývá symetrické lineární rozdělení stavového prostoru. Z hlediska rozlehlosti stavového prostoru není příliš vhodné a je proto vhodné použít druhé možnosti nazvané symetrické nelineární rozdělení stavového prostoru. Toto rozdělení vychází z toho, že není nutné mít stavový prostor rozdělen na velký počet částí v oblasti, které pro nás nejsou příliš důležité. V našem případě v blízkosti námi požadovaných otáček je stavový prostor rozdělen na malé úseky, a čím více se vzdaluje od těchto otáček, tím také velikost těchto úseků roste. Tímto se dá stavový prostor dostatečně minimalizovat.

2. Model učení

Algoritmus Q-učení (Q-Learning) popsal Watkins. Tento algoritmus je velmi jednoduché implementovat. Q-učení je pravděpodobně nejpobulárnější a nejefektivnější RL procedura opakovaně posilovaného učení bez modelu.

Odhad celkového výkonu agenta (Q-funkce)

$$Q_{\pi}(s_i, u) = r(s_i, u) + \sum_{k=i+1}^{\infty} \gamma^{k-i} r(s_k, \pi(s_k)) \quad , \quad 0 < \gamma < 1, \quad \forall s_i, s \in S, \quad \forall u \in U, \quad (1)$$
$$\pi(s) = \arg \max_u Q(s, u)$$

kde	S	konečná množina stavů prostředí,
	U	konečná množina akcí,
	$\pi : S \rightarrow U$	strategie agenta,
	$r : S \times U \rightarrow R$	ohodnocení okamžitého výkonu agenta
	γ	srážkový faktor.

Odhad optimálního celkového výkonu agenta je odhad celkového výkonu agenta při jeho optimální strategii (optimální Q-funkce).

$$\begin{aligned} Q^*(s_i, u) &= \max_{\pi} Q_{\pi}(s_i, u) \\ \pi^*(s) &= \arg \max_u Q^*(s, u) \end{aligned}, \quad \forall s_i, s \in S, \forall u \in U \quad (2)$$

Vhodnost dané akce u ve stavu s je

$$\Delta e(s, u) = \begin{cases} (\gamma\lambda - 1)e(s, u) + 1, & \text{jestliže } (s = s_k) \wedge (u = u_k) \\ (\gamma\lambda - 1)e(s, u), & \text{jinak} \end{cases} \quad \forall s, s_k \in S, \forall u, u_k \in U. \quad (3)$$

Rovnice pro Q-učení tedy je

$$\Delta Q(s, u) = \alpha \left[r(s_k, u_k) + \gamma \max_{u_{k+1}} Q(s_{k+1}, u_{k+1}) - Q(s_k, u_k) \right] e(s, u), \quad \forall s \in S, \forall u \in U. \quad (4)$$

Toto pravidlo zaručuje s pravděpodobností 1 konvergenci Q-hodnot k optimální hodnotě, jestliže je každá akce v každém stavu prostředí vykonána nekonečně mnohokrát během nekonečného počtu kroků a pokud hodnota učícího poměru α vhodně klesá. Pokud Q-hodnoty konvergují k optimální hodnotě, potom je vhodné, aby agent v každém stavu vykonával vždy akci s nejvyšší Q-hodnotou. Q-učení je necitlivé na nastavení, to znamená, že Q-hodnoty budou vždy konvergovat k optimální hodnotě, budou-li všechny dvojice stav-akce vyzkoušeny v dostatečném počtu. Bohužel tato konvergence může být velmi pomalá. Q-učení nepracuje příliš dobře při zobecněných problémech, nebo při stavovém prostoru, jenž je rozsáhlý.

3. Model asynchronního elektromotoru

Jako model byl použit model asynchronního elektromotoru. Řízení otáček je prováděno změnou kmitočtu napájecího proudu. Implementace modelu je provedena následujícími iteračními rovnicemi.

$$\begin{aligned} U_d &= U \cdot \sin(2\pi \cdot f \cdot T) \\ U_q &= -U \cdot \cos(2\pi \cdot f \cdot T) \\ MZ &= 5 \\ D[0] &= I_d \\ D[1] &= I_q \\ I_d &= I_d + (K_0 \cdot T (U_d - R_s \cdot I_d + K_1 \cdot I_D + L_{sr} \cdot W_R \cdot I_Q + K_2 \cdot I_q \cdot W_r)) \\ I_q &= I_q + (K_0 \cdot T (U_q - R_s \cdot I_q + K_1 \cdot I_Q - L_{sr} \cdot W_R \cdot I_D - K_2 \cdot I_d \cdot W_r)) \\ I_D &= I_D + (-K_4 (I_d - D_0) + T \cdot (-K_3 \cdot I_D - W_R \cdot I_Q - K_4 \cdot W_R \cdot I_q)) \\ I_Q &= I_Q + (-K_4 (I_q - D_1) + T \cdot (-K_3 \cdot I_Q + W_R \cdot I_D + K_4 \cdot W_R \cdot I_d)) \\ M_I &= 1.5 \cdot L_{sr} (I_q \cdot I_D - I_Q \cdot I_d) \\ W_R &= W_R + T (M_I - M_Z) / J \\ T &= T + \text{krok} \end{aligned}$$

Kde konstanty K_0 až K_4 jsou vypočteny následujícími vztahy:

$$\begin{aligned} K[0] &= L_r / (L_s \cdot L_r - L_{sr} \cdot L_{sr}) \\ K[1] &= L_{sr} \cdot R_r / L_r \\ K[2] &= L_{sr} \cdot L_{sr} / L_r \\ K[3] &= R_r / L_r \\ K[4] &= L_{sr} / L_r \end{aligned}$$

Proměnné v těchto rovnicích značí a byly nastaveny následovně:

R_s	Odpor statoru	4.37Ω
R_r	Odpor rotoru	2.95Ω
L_s	Indukčnost statoru	0.471 H
L_r	Indukčnost rotoru	0.476 H
L_{sr}	Indukčnost mezi statorem a rotorem	0.459 H
f	Frekvence	50 Hz
U	Napětí	311.13 V
J	Moment setrvačnosti	$0.0015 \text{ Kg}\cdot\text{m}^2$
M_Z	Zátěžný moment	$5 \text{ N}\cdot\text{m}$
Krok	Časový krok	0.00001
U_d	Napětí v ose d	
U_q	Napětí v ose q	
I_d	Proud v ose d	
I_q	Proud v ose q	
I_D	Proud v ose D	
I_Q	Proud v ose Q	
M_I	Interní moment stroje	
W_R	Aktuální otáčky	
T	Čas modelu	

4. Parametry experimentů

Pro všechny testy platí, není-li řečeno jinak, tato pravidla:

V každém experimentu bylo provedeno 5000 pokusů. Úspěšnost byla spočtena každých 250 pokusů jako procentuální poměr úspěšných pokusů k neúspěšným. Za úspěšný pokus byl považován pokus, ve kterém bylo dosaženo požadovaných otáček.

Byly použity tyto akce:

- „přidej frekvenci o 10Hz“
- „přidej frekvenci o 1Hz“
- „udržuj stejnou frekvenci“
- „uber frekvenci o 1Hz“
- „uber frekvenci o 10Hz“.

Počáteční otáčky byly nastaveny náhodně.

Agentova ohodnocovací funkce r byla definována následovně:

Pokud byla akce vhodná, to znamená, že se otáčky modelu přiblížily požadovaným otáčkám, nebo pokud těchto otáček model dosáhl, bylo ohodnocení r funkcí aktuálních otáček v intervalu $(0,1)$. V případě, že akce neměla význam na otáčky modelu, byla hodnota r nastavena na hodnotu 0. Pokud akce nebyla vhodná, to znamená, že se otáčky modelu vzdálily od požadovaných otáček, nebo pokud model dosáhl nulové rychlosti, bylo ohodnocení r funkcí aktuálních otáček v intervalu $(0,-1)$.

Stavový prostor byl nelineární a měl velikost 350. Hodnota 350 udává velikost stavového prostoru vypočtenou jako velikost gridu frekvence*otáčky*akce, což je $7*10*5 = 350$.

Koeficienty pro Q-učení byly nastaveny v následujícím rozsahu:

Učící poměr $\alpha \in (0, 1)$

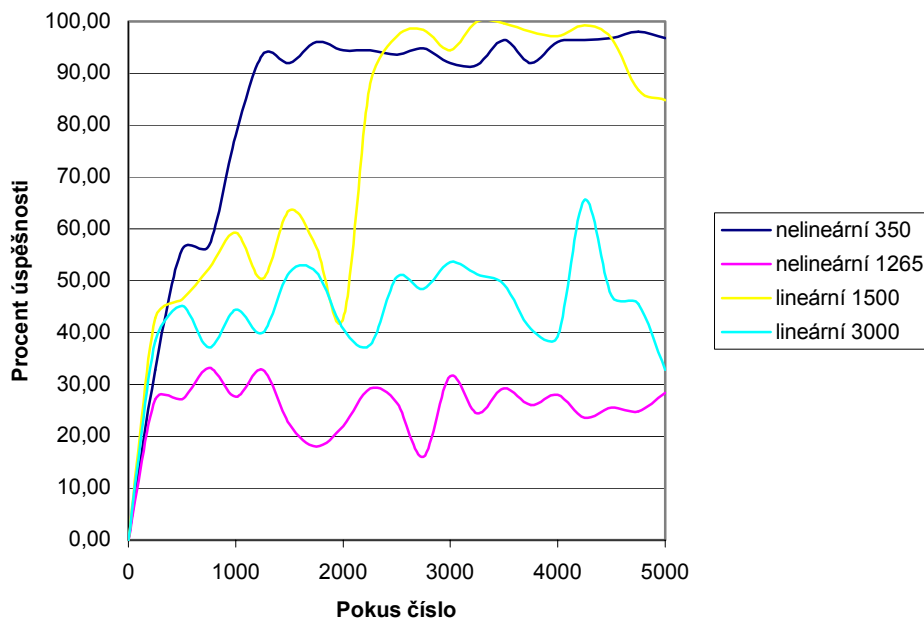
Faktor slevy $\gamma \in (0.7, 1)$

Míra zapomínání $\lambda \in (0.7, 1)$

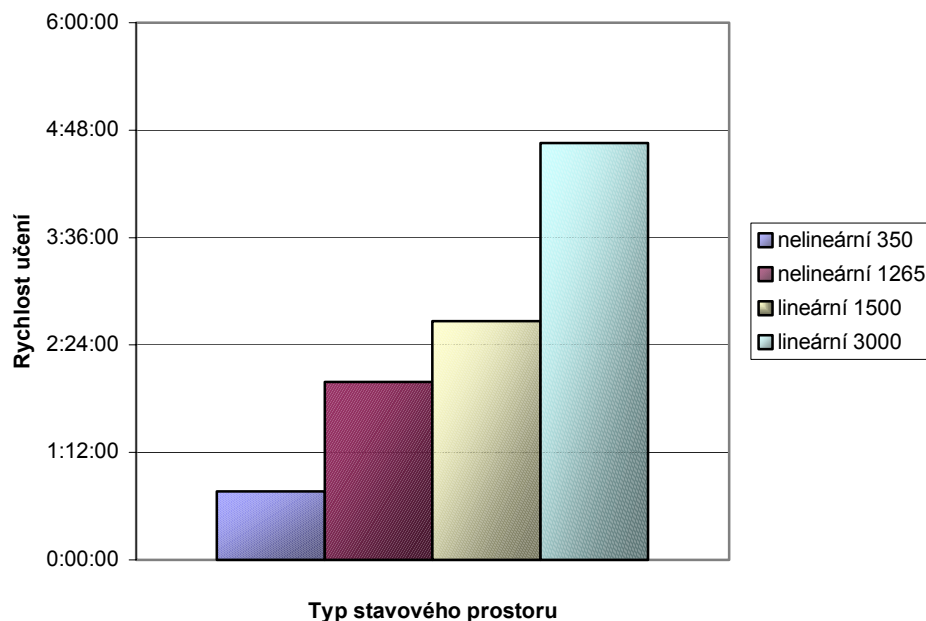
5. Provedené experimenty

5.1. Rozdělení stavového prostoru

Vliv rozdělení stavového prostoru na průběh kvality naučení je zobrazen na obrázku 1. Velikost gridu byla $[7 \times 10 \times 5] = 350$, $[11 \times 23 \times 5] = 1265$, $[10 \times 30 \times 5] = 1500$ a $[15 \times 40 \times 5] = 3000$. Z obrázku je patrné, že po 5000 pokusech se systém nejlépe naučil s nelineárním stavovým prostorem o velikosti 350. Dá se předpokládat, že systém s rozlehlejšími stavovými prostory by se naučil lépe, ale na úkor prodloužení doby výpočtu. Rychlost učení v závislosti na volbě stavového prostoru je na obrázku 2.



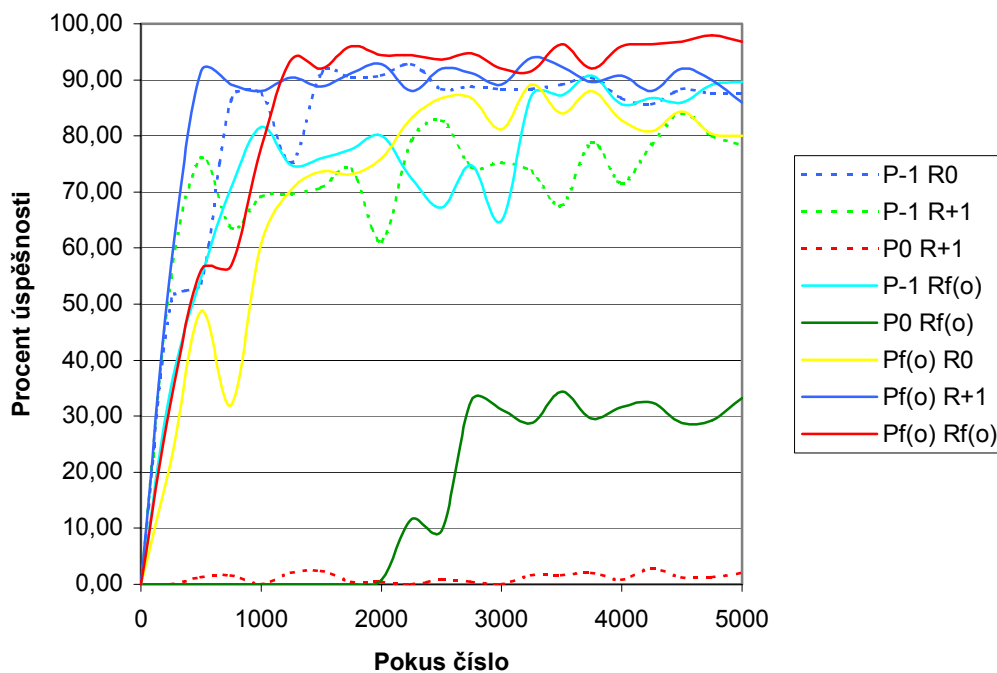
Obr. 1: Vliv velikosti a rozdělení stavového prostoru na průběh kvality naučení



Obr. 2: Rychlost učení v závislosti na volbě stavového prostoru

5.2. Volba odměny

Z obrázku 3 plyne, že ke správnému naučení systému je třeba systém trestat. Různé kombinace odměn a trestů, ať už poměrných či absolutních, přináší podobné výsledky. Ty se pak mohou lišit dle typu řešené úlohy. Problém poměrného trestu a odměny u řídicí úlohy je shodný s problémem poměrného trestu v trestním právu. Každý konkrétní případ je jiný a i když si určíme nějaké měřítko, nikdy správně nepostihneme všechny aspekty. Proto je poměrná odměna i trest nejvýhodnější.

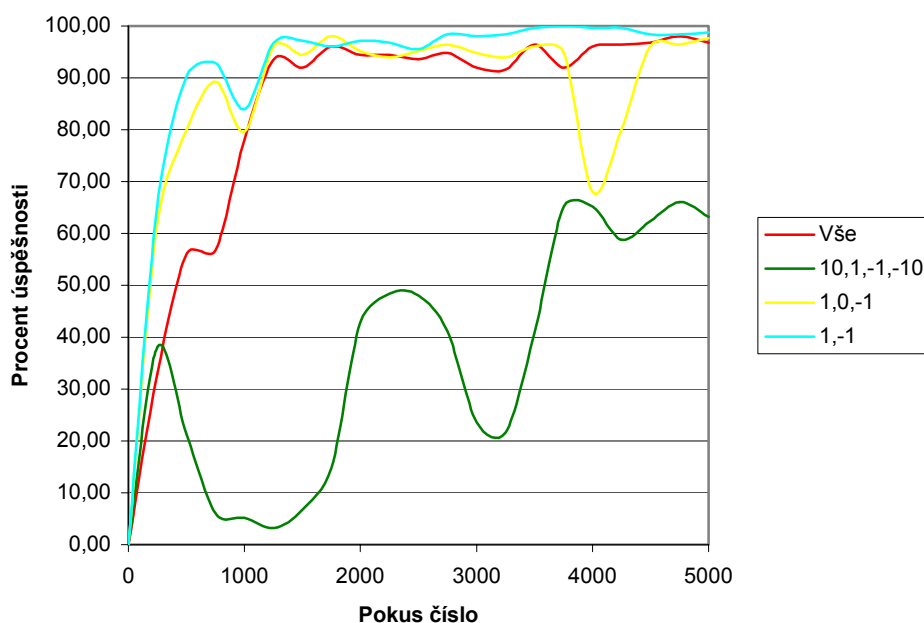


Obr. 3: Graf vlivu různé odměny na průběh kvality naučení

V grafu P udává trest a R odměnu. P-1 a R+1 tudíž znamená, že při špatné akci byl agent potrestán hodnotou $r = -1$ a při vhodné akci byl naopak odměněn hodnotou $r = 1$. Pf(o) udává, že míra trestu byla stanovena na základě aktuálních otáček modelu. Za špatnou akci považujeme tu akci, na základě které se otáčky modelu vzdálí od požadovaných otáček, případně pokud bude dosaženo nulových otáček. Vhodná akce je zase ta, na základě které se aktuální otáčky modelu přiblíží k námi požadovaným otáčkám, případně těchto požadovaných otáček dosáhne.

5.3. Volba akce

Z obrázku 4 je patrné, že nejlepšího naučení se docílilo při použití akcí „přidej frekvenci o 1Hz“ a „uber frekvenci o 1Hz“. V praktických testech se ale ukázalo, že volba těchto akcí není příliš vhodná. Byly proto rozšířeny o akce „přidej frekvenci o 10Hz“, „uber frekvenci o 10Hz“ a „udržuj stejnou frekvenci“, které se osvědčily více.

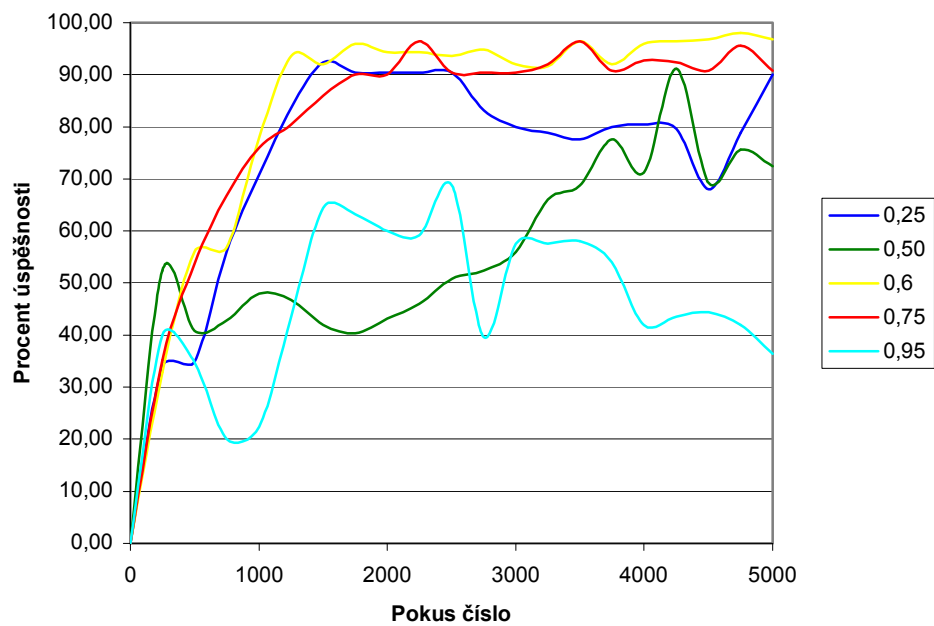


Obr. 4: Graf vlivu různých akcí na průběh kvality naučení

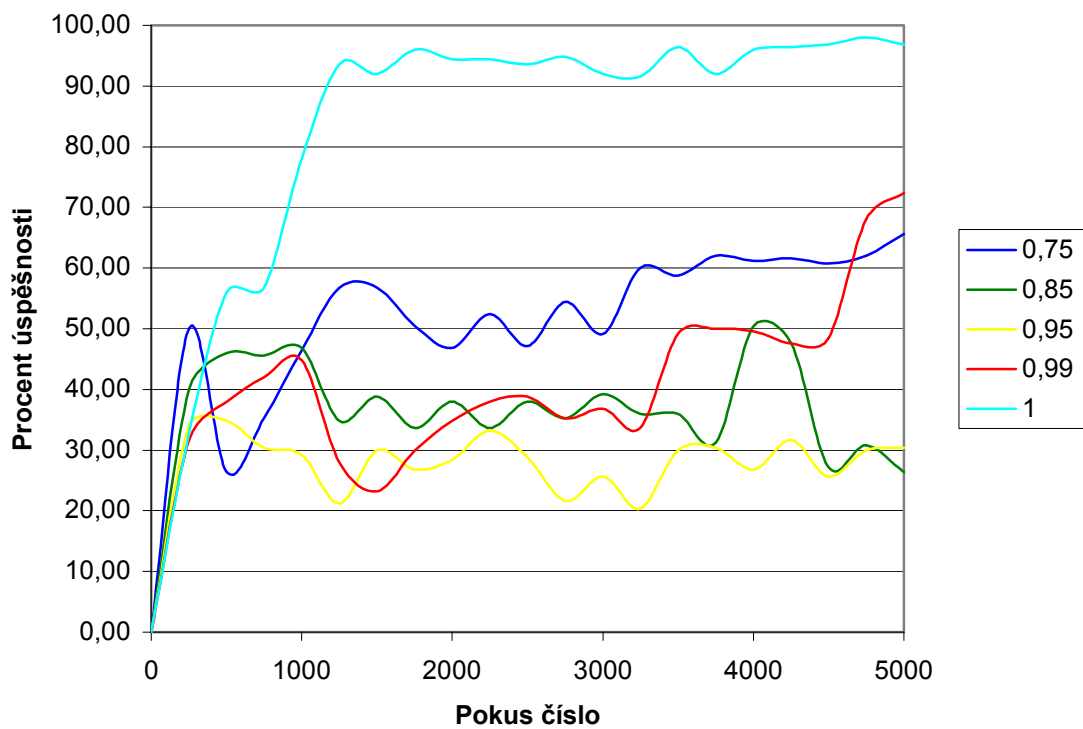
5.4. Volba koeficientů α , γ , λ

Z obrázku 5, 6, 7 je patrné, že nejlepší kvality naučení bylo dosaženo pro následující hodnoty koeficientů Q-učení.

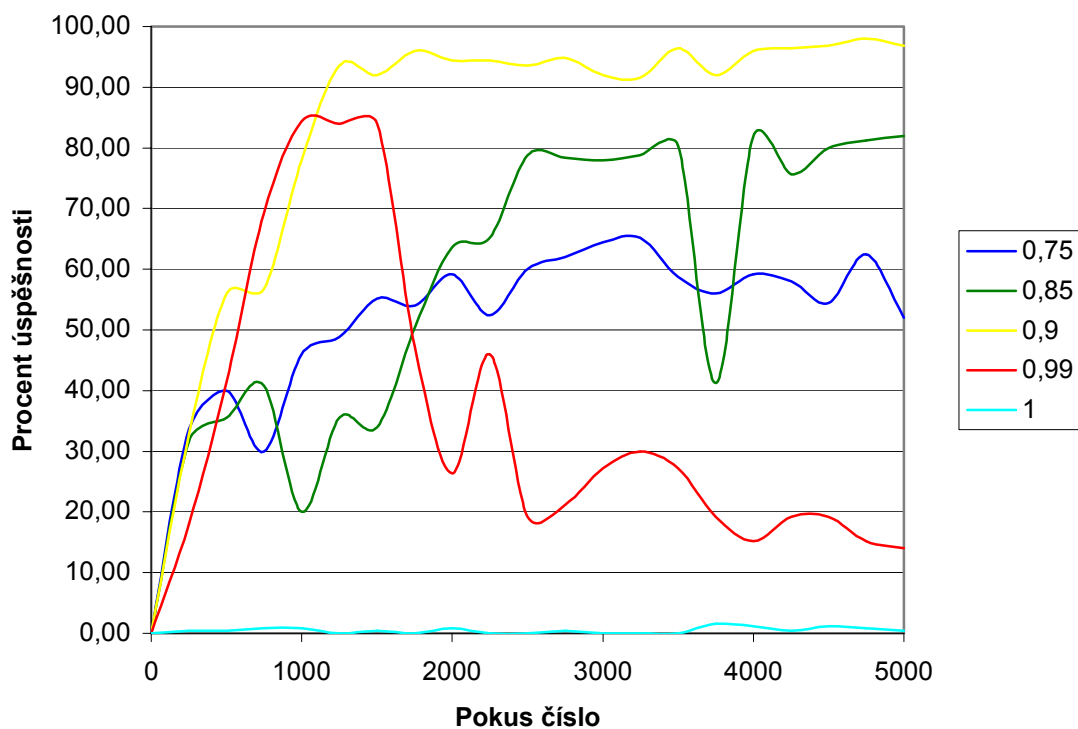
Učící poměr	$\alpha = 0.6$
Srážkový faktor	$\gamma = 1$
Míra zapomínání	$\lambda = 0.9$



Obr. 5: Graf vlivu parametru ALPHA na průběh kvality naučení



Obr. 6: Graf vlivu parametru GAMMA na průběh kvality naučení

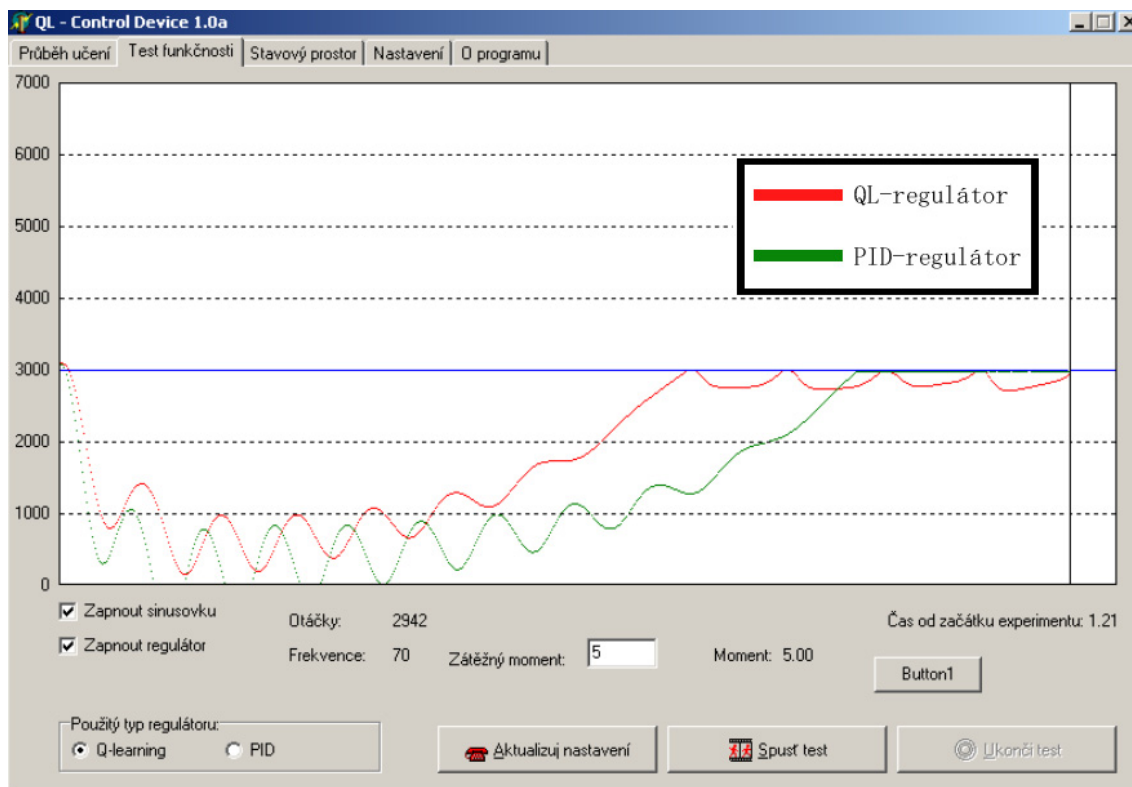


Obr. 7: Graf vlivu parametru LAMBDA na průběh kvality naučení

5.5. Porovnání QL-Regulátoru s PID-Regulátorem

Učení QL-regulátoru se skládalo z učení na dvě úlohy. První úlohou bylo naučit regulátor dosáhnout vhodným způsobem požadovaných otáček. Druhou úlohou bylo tyto požadované otáčky udržet. Systém byl naučen na dosažení požadovaných otáček. PID regulátor byl nastaven pomocí Ziegler Nicholsova pravidla po mírných úpravách na tyto hodnoty.

Zesílení	0.4
Derivační časová konstanta	2
Integrační časová konstanta	50



Obr. 8: Porovnání QL-Regulátoru s PID-Regulátorem

Na obrázku 8 je červenou barvou vyobrazen průběh QL-regulátoru a zelenou barvou průběh PID-regulátoru. Jestliže si rozebereme průběh experimentu zobrazeného na obrázku vidíme, že v první fázi učení dosáhl QL-regulátor lepšího výsledku. Můžeme si také povšimnout toho, že QL-regulátor se choval podobně jako PID-regulátor. Ve druhé fázi učení vidíme, že PID-regulátor dosáhl požadovaných otáček a tuto hodnotu dále udržel. Na rozdíl od něj QL-regulátor tuto hodnotu dosáhl, ale docházelo k periodickým oscilacím. Z toho vyplývá, že ve druhé fázi QL-regulátor nedosáhl stejné kvality jako jeho konkurent PID-regulátor. Tento výsledek byl zapříčiněn tím, že systém nebyl na úlohu udržení otáček naučen.

6. Závěr

Metoda využívá schopnost učení postupně zlepšovat chování asynchronního elektromotoru. Experimentálně bylo vyvozeno, že optimální hodnoty pro nastavení Q-učení jsou tyto:

- učící poměr $\alpha = 0.6$
- srážkový poměr $\gamma = 1$
- faktor zapomínání $\lambda = 0.9$

Nejvýhodnější je poměrná odměna a trest. Jako nejvhodnější akce byly zvoleny akce:

- Dekrementuj řídicí frekvenci o 10Hz
- Dekrementuj řídicí frekvenci o 1Hz

- Udržuj stejnou řídicí frekvenci
- Inkrementuj řídicí frekvenci o 1Hz
- Inkrementuj řídicí frekvenci o 10Hz

Zavedení nelineárního rastru pro všechny veličiny výrazně zlepšuje úspěšnost učení a proto je výhodnější než lineární rastr. Při naučení systému na dosažení požadovaných otáček se systém chová lépe než PID regulátor.

Doporučená literatura

- [1] Richard S. Sutton, Andrew G. Barto. Reinforcement Learning: An Introduction. MIT Press, Cambridge, MA, 1998.
- [2] Leslie P. Kaelbling, Michael L. Littmann, Andrew W. Moore. Reinforcement learning: A survey. Journal of Artificial Intelligence Research, 4, 1996.
- [3] Christopher J.C.H. Watkins and Peter Dayan. Q-learning. Machine Learning, 8(3):279-292, 1992.
- [4] Doc.Ing. Jaroslav Šubrt CSc.: Elektrické stroje, Brno, 1999.
- [5] Doc.Ing. Petr Vavřín CSc.: Teorie automatického řízení 1, SNTL.