

基于情感与环境认知的移动机器人自主导航控制

张惠娣^{1,2}, 刘士荣²

(1. 华东理工大学 自动化研究所, 上海 200237; 2. 杭州电子科技大学 自动化研究所, 浙江 杭州 310018)

摘要: 将基于情感和学习与决策模型引入到基于行为的移动机器人控制体系中, 设计了一种新的自主导航控制系统. 将动力学系统方法用于基本行为设计, 并利用ART2神经网络实现对连续的环境感知状态分类, 将分类结果作为学习与决策算法中的环境认知状态. 通过在线情感和环境认知学习, 形成合理的行为协调机制. 仿真表明, 情感和环境认知能明显地改善学习和决策过程效率, 提高基于行为的移动机器人在未知环境中的自主导航能力.

关键词: 移动机器人; 情感; 认知; ART2神经网络; 行为协调; 自主导航

中图分类号: TP242 **文献标识码:** A

Autonomous navigation control for mobile robots based on emotion and environment cognition

ZHANG Hui-di^{1,2}, LIU Shi-rong²

(1. Research Institute of Automation, East China University of Science and Technology, Shanghai 200237, China;

2. Institute of Automation, Hangzhou Dianzi University, Hangzhou Zhejiang 310018, China)

Abstract: A new autonomous navigation control system for mobile robots is presented, in which the learning and decision making model based on emotion and cognition is introduced into the control architecture of behavior-based robot system. Dynamic system approach is applied to design the basic behaviors. Continuous environment perceptions are classified into categories by Adaptive Resonance Theory-2 (ART2) networks, which are used as the environmental cognitive states in learning and decision making. Rational behavior-coordination mechanism is developed by the on-line learning of emotion and environmental cognitive. Simulations have demonstrated that the emotion and environmental cognition can obviously improve the efficiency of the learning and decision making process, and enhance the capability of autonomous navigation in behavior-based robot system under unknown environment.

Key words: mobile robot; emotion; cognition; ART2 networks; behavior coordination; autonomous navigation

1 引言(Introduction)

已有的拟人控制理论主要是维纳的反馈控制论和人工智能^[1], 与人脑控制模式还有很大差别. 人脑控制模式是: 感知觉+情感决定行为. 神经科学、心理学和认知科学等领域的研究表明: 情感在推理、学习、记忆、决策、智能行为等方面扮演着至关重要的角色^[2]. 情感在决策中的作用模式的机器实现, 主要是模拟人脑的控制模式, 建立“感知觉+情感决定行为”(人脑控制模式)的数学模型^[3]. 情感智能的研究正成为人工智能领域的一个重要研究方向, 新的相关研究内容和应用成果正在不断产生^[4~6].

受人脑控制模式启发, MIT情感计算研究小组的Ahn和Picard^[7,8]提出了一个新的强化学习模型: 基于情感和学习与决策模型. 本文将该模型融入到基于行为的机器人控制体系中, 采用ART2神经网络对环境状态进行分类, 并利用动力学系统方法实现基本行为设计. 有机地将情感、认知、行为学习和控制结合起来, 建立了一种基于情感和环境认知的行为学习和决策的移动机器人导航控制系统.

2 自主导航控制系统结构(Architecture of autonomous navigation control system)

Ahn和Picard^[7,8]所提出的基于情感认知的学习与决策算法同时考虑了来自情感的内在奖励和来自

为动力学系统的稳定性, 必须满足 $|\dot{\phi}| \gg |\dot{\psi}_i|$ ($i = \text{goal, obs, or wall}$). 因此相关的行为参数必须满足以下条件^[11]:

$$k_i \gg \lambda_{i,\phi} \text{ 且 } k_i \gg \lambda_{i,\nu}, i = \text{goal, obs, or wall}. \quad (8)$$

与常规的移动机器人线速度和角速度的设定值为常数的情况相比, 基于动力学系统方法的行为设计能使移动机器人根据工作环境状态自主地调整导航方向和导航速度, 从而可以确保机器人实现更为安全的导航控制。

4 基于ART2神经网络局部环境感知(Perception of local environment based on ART2 networks)

环境状态就是感知模型的结果, 如果直接用传感器的输出值来确定, 会导致环境状态的维数太高, 不利于学习. 于是可根据移动机器人当前位置的左右两侧以及前进方向障碍物的趋近度来确定环境状态。

因为不同方向的障碍物对导航控制任务具有不同的影响, 所以在趋近度定义中考虑了方向角的作用^[12]. 根据机器人的传感器分布(见图2), 左右两侧和前进方向障碍物的趋近度分别定义为

$$M_R(t) = \sum_{i=1}^3 \frac{|\cos \alpha_i|}{dS_i + m_0}, \quad (9)$$

$$M_L(t) = \sum_{i=6}^8 \frac{|\cos \alpha_i|}{dS_i + m_0}, \quad (10)$$

$$M_F(t) = \sum_{i=4}^5 \frac{|\cos \alpha_i|}{dS_i + m_0}. \quad (11)$$

其中: 下标R, L, F分别表示机器人右边、左边、前向; α_i 是机器人局部坐标系中的第*i*个传感器方向角; dS_i 是第*i*个传感器输出; $m_0 > 0$ 是裕量常数。

$[M_R(t), M_F(t), M_L(t)]$ 表示了机器人所处的周围环境状态. 机器人所处的环境状态是连续的时间序列, 行为空间和环境状态空间的组合将会引起组合爆炸问题. 为此, 本文用ART2网络^[13]对环境状态的时间序列进行分类, 以获得机器人对环境的认知状态. ART2网络使相似的环境归属于同一个认知状态, 从而可大大地减小环境认知状态空间。

机器人在导航过程中, 根据公式(9)~(11)将传感器输出值转换为趋近度, 从而获得环境状态时间序列. 随后将该时间序列作为ART2神经网络的输入进行分类, 或获得1个新的环境认知状态。

假设机器人在每一时刻只有1个认知状态. $f(c^t, d^t, c^{t+1})$ 表示执行决策 d^t 后, 从认知状态 c^t 转移到认知状态 c^{t+1} 的次数, 即

$$f(c^t, d^t, c^{t+1}) \leftarrow f(c^t, d^t, c^{t+1}) + 1. \quad (12)$$

认知状态传递模型 $P(c^{t+1}|c^t, d^t)$ 表示从当前认知状态到下一认知状态的转移过程, 可用概率分布来描述:

$$P(c'|c^t, d^t) = \frac{f(c', d^t, c^t)}{\sum_{k=1}^{|C|} f(c^t, d^t, k)}, \quad c' \in \mathbb{C}. \quad (13)$$

5 自主导航控制算法(Autonomous navigation control algorithm)

5.1 情感模型与内在奖励(Affective model and intrinsic reward)

人类行为的选择往往是以追求正向情感最大化、负向情感最小化为目标. 因此, 为了简化人工情感系统, 采用概率分布来描述情感模型, 即用概率分布来表示每个情感分量. 设

$$a_e = (a_{e,1}, a_{e,2}), e \in A = \{1, \dots, |A|\}. \quad (14)$$

其中 $a_{e,1}, a_{e,2}$ 分别表示感觉坏(feeling bad)与感觉好(feeling good)的条件概率, 并且 $a_{e,1} + a_{e,2} = 1$. 机器人的情感状态则可用多个条件概率对来表示:

$$a = a_e = (a_1, \dots, a_{|A|}). \quad (15)$$

本文以移动机器人从初始位置出发移动到指定的目标位置的基本任务为例来研究导航控制策略, 因此只考虑最基本的want情感, 即 $a = (a_{1,1}, a_{1,2})$ 。

当决策完成后, 被选中的行为作用于环境与机器人本身, 引起环境状态和机器人认知状态改变. 此时机器人的情感系统重新评估当前的情况:

$$\delta_E = R_{\text{ext}}(c^{t+1}|c^t, d^t) - \max_{c'} \left(\sum_{c'} R_{\text{ext}}(c'|c^t, d^t) P(c'|c^t, d^t) \right). \quad (16)$$

当 $\delta_E > 0$ 时, 表示在相同条件下(即相同决策和认知状态)情感的下一状态为feeling good的概率较大, 反之亦然。

根据评估值产生新的情感状态^[8]:

$$a_{1,v_e^{t+1}}^{t+1} = P(v_e^{t+1}|c^{t+1}, c^t, a^t, d^t). \quad (17)$$

其中:

$$v_e^{t+1} = \begin{cases} 1, & \text{feeling bad,} \\ 2, & \text{feeling good,} \end{cases}$$

$$P(v_e = 1|c^{t+1}, c^t, a^t, d^t) = P(v_e = 1|c^{t+1}, c^t, a^t, d^t) - \varepsilon \delta_E, \quad (18)$$

$$P(v_e = 2|c^{t+1}, c^t, a^t, d^t) = P(v_e = 2|c^{t+1}, c^t, a^t, d^t) + \varepsilon \delta_E. \quad (19)$$

式(16)~(19)表明环境的一些重要变化会引起情感状态的变化, 也就是说情感状态是对全局环境的一

个概括反映. 因此, 利用情感模块产生学习和决策的内在奖励信号是合理的. 根据所设计的情感模型, 来自情感的奖励信号定义为

$$R(v_e^{t+1}) = \begin{cases} 1, & v_e^{t+1} = 1, \text{ feeling bad,} \\ -1, & v_e^{t+1} = 2, \text{ feeling good.} \end{cases} \quad (20)$$

5.2 外部奖励(Extrinsic reward)

移动机器人自主导航任务是一个多目标行为, 即避开障碍物和趋近目的地. 因此, 外部奖励信号应该包括两部分:

1) 路径的安全性.

为了防止机器人的某一部位与障碍物发生碰撞, 应尽量使其与障碍物保持一定安全距离. 机器人与周围障碍物的趋近度定义为

$$M_c(t) = \sum_{i=1}^8 \frac{|\cos \alpha_i|}{dS_i + m_0}. \quad (21)$$

$M_c(t)$ 体现了机器人距障碍物的综合相对位置关系, 因此取奖励信号为

$$r_{\text{ext}1}(t) = -(M_c(t+1) - M_c(t)). \quad (22)$$

2) 趋近目标.

在机器人执行趋近目标的行为时, 需要考虑机器人与目标的距离和角度关系, 令奖励信号为

$$r_{\text{ext}2}(t) = -p_1 \cdot \Delta D_g(t) - p_2 \cdot \Delta \varphi_g(t). \quad (23)$$

其中: ΔD_g 表示机器人与目标距离的变化量, $\Delta \varphi_g$ 是机器人与目标夹角的变化量, p_1, p_2 均为非负参数.

机器人与障碍物的距离不同导致其行为的侧重点也应该不同, 取外部奖励值为

$$r_{\text{ext}}(t) = \begin{cases} w_{r2} \cdot r_{\text{ext}2}(t), & M_c(t+1) < M_{\text{cmin}}, \\ -C_{\text{ext}}, & M_c(t+1) > M_{\text{cmax}}, \\ w_{r1} \cdot r_{\text{ext}1}(t) + w_{r2} \cdot r_{\text{ext}2}(t), & \text{其他.} \end{cases} \quad (24)$$

其中: w_{r1}, w_{r2} 分别为权重系数, M_{cmax} 为所允许的最大趋近度, M_{cmin} 为安全趋近度. $C_{\text{ext}} > 0$ 为一个较大的惩罚系数.

定义外部奖励模型为

$$R_{\text{ext}}(c^{t+1}|c^t, d^t) = (1 - \gamma) \cdot R_{\text{ext}}(c^{t+1}|c^t, d^t) + \gamma \cdot r_{\text{ext}}(t). \quad (25)$$

其中学习率 $\gamma = 1/f(c^t, d^t, c^{t+1})$. 外部奖励模型 $R_{\text{ext}}(c'|c, d)$ 是基于即时外部奖励信号 r_{ext} 的一个动态模型, 它反映了学习过程中, 移动机器人在决策 d 作用下从认知状态 c 转移到认知状态 c' 时所获得的综合外部奖励.

5.3 学习与决策算法(Learning and decision-making algorithm)

具体的学习和决策步骤如下:

第1步 在当前的认知状态 c^t 、情感状态 a^t 下, 根据 Boltzmann 分布做出行为策略 d^t :

$$Pr(d^t|c^t, a^t) = \frac{\exp(Q_{DM}(c^t, a^t, d^t)/T)}{\sum_{d=1}^D \exp(Q_{DM}(c^t, a^t, d)/T)}. \quad (26)$$

其中: T 为“温度”系数, 决策值 Q_{DM} 由来自认知评价系统的外部决策值 Q_{ext} 和来自情感模型的内在决策值 Q_{int} 组成:

$$Q_{DM}(c^t, a^t, d) = Q_{\text{ext}}(c^t, d) + \eta Q_{\text{int}}(c^t, a^t, d). \quad (27)$$

其中参数 η 表示来自情感模型的奖励信号对整个系统决策的影响.

第2步 执行决策 d^t , 由 ART2 神经网络获得 1 个新的环境认知状态 c^{t+1} , 并根据式(24)得到外部奖励 r_{ext} .

第3步 利用 $\langle c^t, d^t, r_{\text{ext}}, c^{t+1} \rangle$ 更新认知状态传递模型 $P(c^{t+1}|c^t, d^t)$ 、外部奖励模型 $R_{\text{ext}}(c^{t+1}|c^t, d^t)$.

第4步 更新情感状态传递模型 $P(v_e^{t+1}|c^{t+1}, c^t, a^t, d^t)$.

第5步 对所有 $c \in C, a \in A, d \in D$ 分别更新外部决策值 Q_{ext} 和内在决策值 Q_{int} :

$$Q_{\text{ext}}(c, d) = \sum_{c'=1}^{|C|} R_{\text{ext}}(c'|c, d)P(c'|c, d) + \alpha \sum_{c'=1}^{|C|} \max_{d'} Q_{\text{ext}}(c', d')P(c'|c, d), \quad (28)$$

$$Q_{\text{int}}(c, a, d) = \sum_{c'=1}^{|C|} \sum_{v_e=1}^2 R_e(v_e)P(v_e|c', c, a, d)P(c'|c, d). \quad (29)$$

第6步 由式(17)~(19)获得一个新的情感状态 $a^{t+1} = (a_{1,1}^{t+1}, a_{1,2}^{t+1})$.

第7步 $t \leftarrow t + 1$, 返回第1步.

上述算法中通过对行为执行结果的评价获得新的外部奖励信号和内部奖励信号, 并以此来更新外部决策值和内在决策值. 引入来自情感的内在奖励信号可提高机器人的学习和决策速度, 并且随着参数 η 的增大, 来自情感的奖励信号对整个系统决策的影响也随之增强.

6 仿真研究(Simulation study)

为了说明所提出的移动机器人导航控制策略的有效性, 进行了仿真实验研究. 定义机器人每趋近障

碍物一次, 完成一次学习过程. 当机器人与障碍物接近时, 回到起点重新开始学习, 即在上次学习结果的基础上重新进行决策值的调整. 从实验来看, 经过几次学习, 机器人就学会在复杂环境中行走.

实验中, 情感对决策的影响因子 η 分别取:

Case 1 $\eta = 0$, 采用ART2环境认知模型, 仅利用外在奖励信号, 即基于环境认知的学习与决策过程.

Case 2 $\eta = 1$, 采用ART2环境认知模型, 同时利用外在和内在的奖励信号, 并且在决策过程中情感认知和环境认知具有相同的重要性.

Case 3 $\eta = 5$, 采用 ART2 环境认知模型, 与 Case 2 相比, 更加强调了来自情感系统的奖励信号对系统决策的影响.

Case 4 $\eta = 5$, 采用常规环境模型, 即直接采用趋近度(式(9)~(11))表示的环境模型, 但强调了来自情感系统的奖励信号对系统决策的影响.

图3~图6分别给出了前述4种情况的自主导航学习过程. A,B分别为起点和目标点. 当没有情感奖励信号时, 需要3次学习后机器人才能自主地实现从起点到目标点的导航控制, 如图3所示. 当引入来自情感模型的奖励信号后, $\eta = 1$ 时只要两次学习过程后就能实现自主导航任务, 如图4所示. 而由图5可知, $\eta = 5$ 时在学习过程中就可以完成自主导航任务. 由此可见, 在基于环境认知的学习过程中引入情感模型能明显提高学习速度. 图5和图6的结果表明: 由于ART2神经网络能够将相似的环境状态分类到同一模式类, 所以有效地减少相似状态下的决策时间, 从而提高了导航效率.

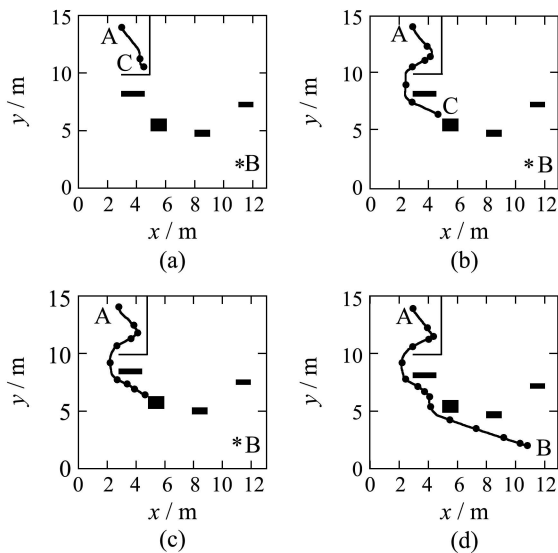


图3 自主导航学习过程(Case 1)

Fig. 3 Learning process of autonomous navigation of Case 1

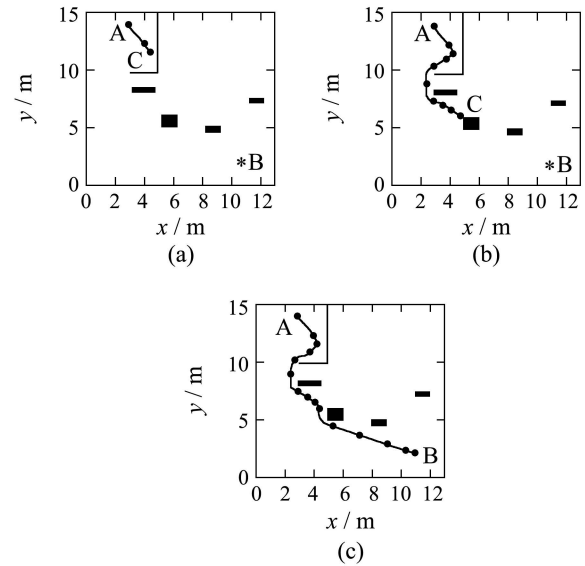


图4 自主导航学习过程(Case 2)

Fig. 4 Learning process of autonomous navigation of Case 2

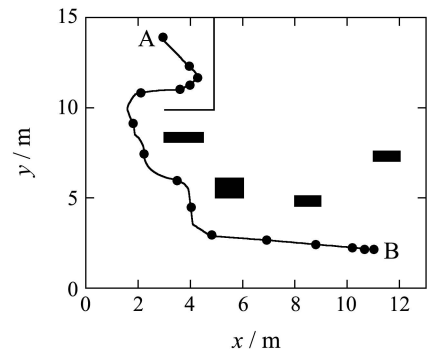


图5 自主导航学习过程(Case 3)

Fig. 5 Learning process of autonomous navigation of Case 3

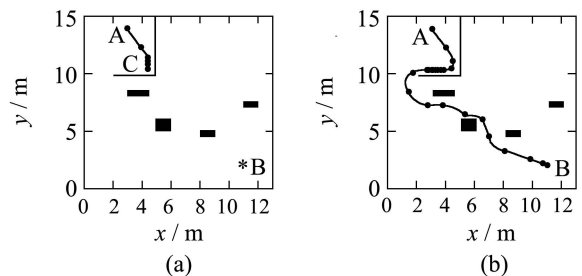


图6 自主导航学习过程(Case 4)

Fig. 6 Learning process of autonomous navigation of Case 4

从图3.4和图6的学习过程中, 可以发现机器人都是在陷入局部极小点处(图中所示的C点处)开始重新学习, 这表明本文所提出基于情感认知学习的行为协调策略能够用于解决机器人导航过程中的局部极小问题, 从而可以提高机器人的自主导航性能.

表1给出了机器人学习过程所需的运行步数和路径长度, 表明机器人的运行步数和路径长度都会随着 η 的增大而减少. 可见在基于环境认知的学习过程中引入情感模型能明显提高学习速度.

表1 4种情况下的仿真结果比较

Table 1 Comparison of simulation results of 4 cases

Case	运行步数	路径长度/m	学习次数
1	350	43.21	4
2	294	33.04	3
3	157	20.43	1
4	260	25.71	2

7 结论(Conclusion)

在基于情感和环境认知的移动机器人自主导航控制系统中, 根据情感和认知的综合评估结果, 分别产生来自情感的内部奖励和来自认知的外部奖励, 以加强合适的行为所对应的权值, 削弱不合适的行为对应的权值, 从而在线地形成了合理的行为协调机制. 仿真结果表明通过将情感与认知结合能够改善学习速度, 提高机器人在未知环境中的自主导航能力.

参考文献(References):

- [1] 国家自然科学基金委员会. 自然科学学科发展战略调研报告: 自动化科学与技术[M]. 北京: 科学出版社, 1995.
(National Natural Science Foundation of China. *Study on Development Strategy of Natural Science: Automation Science and Technology*[M]. Beijing: Science Press, 1995.)
- [2] PICARD R W. *Affective Computing*[M]. Cambridge, Mass, London, England: MIT Press, 1997.
- [3] 王志良. 人工心理学—关于更接近人脑工作模式的科学[J]. 北京科技大学学报, 2000, 22(5): 479–481.
(WANG Zhiliang. Artificial psychology-A most accessible science research to human brain[J]. *Journal of University of Science and Technology Beijing*, 2000, 22(5): 479–481.)
- [4] MOREN J, BALKENIUS. Reflections on emotion[J]. *Cybernetics and Systems*, 2000, 30(6): 699–703.
- [5] VELASQUEZ J D. An Emotion-Based Approach to Robotics[C]//*Proceedings of the 1999 International Conference on Intelligent Robots and Systems*. Piscataway: IEEE Press, 1999: 235–240.
- [6] GADANHO S C, HALLAM J. Emotion-triggered learning in autonomous robot control[J]. *Cybernetics and Systems*, 2001, 32(5): 531–559.
- [7] AHN H, PICARD R W. Affective-cognitive learning and decision making: A motivational reward framework for affective agents[C]//*Proceedings of the 1st International Conference on Affective Computing and Intelligent Interaction*. Beijing, China: Springer, 2005, 866–873.
- [8] AHN H, PICARD R W. Affective-cognitive learning and decision making: The role of emotions[C]//*Proceedings of the 18th European Meeting on Cybernetics and Systems Research*. Vienna, Austria: Austrian Society for Cybernetics Studies, 2006.
- [9] SCHONER G, DOSE M. A dynamical systems approach to task-level system integration used to plan and control autonomous vehicle motion[J]. *Robotics and Autonomous Systems*, 1992, 10(4): 253–267.
- [10] BICHO E, SCHONER G. The dynamic approach to autonomous robotics demonstrated on a low-level vehicle platform[J]. *Robotics and Autonomous Systems*, 1997, 21(1): 23–35.
- [11] ALTHAUS P. *Indoor navigation for mobile robots: control and representations*[D]. Stockholm: Royal Institute of Technology, 2003.
- [12] 谭民, 王硕, 曹志强. 多机器人系统[M]. 北京: 清华大学出版社, 2005.
(TAN Min, WANG Shuo, CAO Zhiqiang. *Multi-robot Systems*[M]. Beijing: Tsinghua University Press, 2005.)
- [13] CARPENTER G A, GROSSBERG S. ART-2: self-organization of stable category recognition codes for analog input pattern[J]. *Applied Optics*, 1987, 26(23): 4919–4930.

作者简介:

张惠娣 (1974—), 女, 博士研究生, 研究方向为智能机器人系统、非线性系统建模与控制, E-mail: nbzhd@hotmail.com;

刘士荣 (1952—), 男, 教授, 博士生导师, 研究方向为复杂过程模型化、控制与优化、智能机器人与智能控制等, E-mail: liushi-rong@hdu.edu.cn.