

به نام پروردگار دانایی

نام و نام خانوادگی: امیرمسعود فرهمند

عنوان پایان نامه: یادگیری و تکامل در معماری های رفتارگرایی سلسله مراتبی

استاد راهنما: دکتر مجید نیلی احمدآبادی

استادان مشاور: دکتر کارو لوکس - دکتر بابک اعرابی

اینجانب امیرمسعود فرهمند نویسنده پایان نامه کارشناسی ارشد در گروه مهندسی برق و کامپیوتر گرایش کنترل بدین وسیله اعلام می نمایم که توسط اساتید راهنما و سمینار در جریان قوانین کپی رایت و روش درست اقتباس و نقل مطالب از سایر مراجع و مواخذ قرار گرفته ام و متعهد به حفظ امانت داری و قدردانی از زحمات سایر محققین و نویسندگان می باشم. بدین وسیله اعلام می نمایم که مسوولیت کلیه مطالب درج شده با اینجانب می باشد و در صورت استفاده از اشکال، جداول و مطالب از سایر منابع بلافاصله مرجع آن ذکر شده است و سایر مطالب از کار تحقیقی اینجانب استخراج شده است و امانت داری را به صورت کامل رعایت نموده ام. در صورتی که خلاف این مطلب ثابت شود، مسوولیت کلیه عواقب قانونی با شخص اینجانب است.

امیرمسعود فرهمند



دانشگاه تهران

دانشکده فنی
گروه مهندسی برق و کامپیوتر

عنوان:

یادگیری و تکامل در معماری‌های رفتارگرای سلسله‌مراتبی

نگارش:

امیرمسعود فرهنگمند

استاد راهنما:

دکتر مجید نیلی احمدآبادی

استادان مشاور:

دکتر کارو لوکس

دکتر بابک اعرابی

پایان‌نامه برای دریافت درجه کارشناسی ارشد در رشته

مهندسی برق - گرایش کنترل

امرداد ۱۳۸۴

دانشگاه تهران
دانشکده فنی
گروه مهندسی برق و کامپیوتر

عنوان:

یادگیری و تکامل در معماری‌های رفتارگرای سلسله‌مراتبی

نگارش:

امیرمسعود فرهمند

پایان‌نامه برای دریافت درجه کارشناسی ارشد در رشته
مهندسی برق - گرایش کنترل

از این پایان‌نامه در تاریخ ۱۳۸۴/۵/۵ در مقابل هیات داوران دفاع به عمل آمد و مورد
تصویب قرار گرفت.

دکتر جواد فیض	معاونت تحصیلات تکمیلی دانشکده فنی
دکتر پرویز جبه‌دار مارالانی	مدیر گروه آموزشی
دکتر سید مهدی فخرایی	سرپرست تحصیلات تکمیلی گروه
دکتر مجید نیلی احمدآبادی	استاد راهنما
دکتر کارو لوکس	استاد مشاور ۱
دکتر بابک نجار اعرابی	استاد مشاور ۲
دکتر بهزاد مشیری	عضو هیات داوران
دکتر علیرضا فاتحی	عضو هیات داوران

چکیده

هدف اصلی این پژوهش، خودکارسازی فرآیند طراحی عامل‌های هوشمند رفتارگرایی سلسله‌مراتبی است. برای انجام این کار فرض می‌شود که سیگنال تقویت‌ای توسط طراح مشخص شده که بیان‌گر عمل‌کرد عامل است. روش‌های پیشنهادی در این گزارش سعی در بهبود آن معیار دارند. در این پایان‌نامه مساله‌ی طراحی عامل‌هایی با معماری رفتارگرایی سلسله‌مراتبی به دو زیرمساله‌ی طراحی رفتارهای معماری و طراحی ساختار معماری تقسیم شده است که حل هم‌زمان این دو منجر به حل مساله‌ی طراحی معماری می‌شود. در این پژوهش دو روی‌کرد مبتنی بر یادگیری و مبتنی بر تکامل به عنوان روش‌های خودکارسازی فرآیند طراحی معرفی شده است.

با استفاده از ایده‌های مبتنی بر یادگیری، هر دو مساله‌ی یافتن ساختار و رفتار مطلوب مطرح شده و سپس پاسخ مناسب برای حل آن‌ها ارائه می‌شود. برای حل مساله‌ی یادگیری ساختار، دو بازنمایی مختلف دانش موسوم به بازنمایی مرتبه صفر و مرتبه یک معرفی می‌شود که با تجزیه‌ی ارزش کل عامل به اجزای ساده‌تر، مسایل دشواری چون تقسیم امتیاز و به‌روزرسانی دانش را قابل حل می‌کند. ایده‌های یادگیری برای ایجاد نگاشت مناسب هر رفتار نیز به کار رفته است و با معرفی عمل‌ای مجازی، اجازه‌ی همکاری بین رفتارها داده شده است.

تکامل هم‌کارانه‌ی رفتارها نیز به عنوان روش‌ای برای تولید رفتارهای جدید استفاده شده است. در این پژوهش برخلاف اکثر پژوهش‌های پیشین، رفتارها به طور مجزا تکامل می‌یابند و سپس با استفاده از یادگیری ساختار به نحوی در معماری قرار می‌گیرند که به بیش‌ترین بازده ممکن منجر شوند.

روش‌های الهام‌گرفته از فرهنگ به نام الگوریتم‌های ممیک نیز معرفی شده‌اند که باعث بالا رفتن بازده استفاده‌ی هم‌زمان از یادگیری و تکامل می‌شود.

روش‌های پیش‌نهادی بر دو بستر بلندکردن جسم توسط چند روبات و مساله‌ای مجرد آزموده شده‌اند و نتایج قابل رقابت با طراحی‌های انسان به دست آمده است.

در این پژوهش علاوه بر نتایج اصلی مربوط به طراحی‌ی خودکار معماری‌ی رفتارگرایی سلسله‌مراتبی، نتایج فرعی‌ای چون تحلیل احتمالاتی‌ی معماری‌ی رفتارگرایی سلسله‌مراتبی‌ی به‌کار رفته در این گزارش و همچنین تاثیر عدم قطعیت سیگنال تقویت بر تابع ارزش و تابع سیاست عامل نیز بررسی شده است.

تشکر

این پایان نامه نوشته نمی شد مگر با کمک استادان و دوستان عزیزم که در این مدت مرا همه گونه یاری کردند.

تجربه ی کار مشترک با مجید نیلی یکی از بهترین تجربه های علمی من بوده است. او نه تنها افق های تازه ای از هوش مصنوعی را به روی ام گشود، بلکه شیوه ی کار آکادمیک را نیز به من یاد داد. او مرا در انتخاب مسیر پژوهش آزاد گذاشت و در عوض به عنوان نقادی بسیار خوب در مورد ایده های ام - چه آن هایی که جا افتاده بودند و چه آن هایی که تنها جرقه ای بوده اند و من یک دفتر پر از آن ها ثبت کرده ام - نظر داد. ما با وجود آن که در بعضی از زمینه های علمی با هم اختلاف نظر داشتیم اما به خوبی در بسیاری از موارد به تفاهم رسیدیم و با کم ترین تقابل ای پروژه را به پیش بردیم. علاوه بر این ساعت ها وقتی که با هم صرف تصحیح مقالات کردیم جزو تجربه های نادری بود که تنها از یک استاد راهنمای ایده آل بر می آید.

کارو لوکس یکی از بهترین استادانی بوده است که تاکنون داشته ام. او مخزن دانش ای است که هر کس به اندازه ی توان خود از آن برداشت می کند. هر بار که با او صحبت می کردم یا سر کلاس های اش می نشستم، انبوه ای از ایده های جدید به ذهن ام خطور می کرد. او بدون شک یکی از شکل دهنده های اصلی اندک دانش ای است که من در این مدت کسب کرده ام.

بابک اعرابی، دیگر استاد مشاورم، با اشاره های مفیدش درک مرا نسبت به آن چه انجام می دهم و ارتباطش با دیگر حوزه های مربوط -مانند شناسایی الگو و یا کنترل- بالا برد. علاوه بر این، برخورد مثبت و امیدبخش او همیشه برای ام ارزش مند بوده و مرا به ادامه ی کار -چه به پروژه ام مربوط بوده باشد و چه به کارهای دیگر- امیدوار می کرده است.

در این دوران سه ساله دوستان خوبی داشته ام. دوستانی که وقت نداشتن ها و بداخلاقی های یک دانش جوی فوق لیسانس را تحمل کرده و او را فراموش نکردند. دوستانی که ساعت ها با من سر موضوعات علمی، فلسفی و یا هر چیز دیگری که بتوان تصور کرد بحث می کردند و ساعات هیجان انگیز

و لذت بخشی را برای ام به وجود آوردند. نمی توانم همه شان را نام ببرم و بگویم چرا از بابت وجودشان خوش حال ام. هر کدام دلیل خاص خود را دارند که با یکی دو کلمه خلاصه پذیر نیست. در این جا تنها نام آن هایی را می آورم که در این مدت به طور پیوسته با آن ها ارتباط داشته ام. از احمد ایمانی پور، احمد پورصابری، آیدین کروی، بتسابه درودیان، پریسا درویش زاده، پریسا رشیدی، پریسا میرشمس، پویان شیوا، پیمان حامدحسینی، حسین مباحی، رامین مهران، رضا پدرامی، روزبه دانشور، روزبه فضل، سعید بیتی، محسن محمدی تاکامی، محمد سلیمانی، محمد شاه محمدی، محمد قشلاقی، محمد میلاسی، مریم گودرزی، هاجر همایی و خیلی های دیگر که نام شان را نیاورده ام متشکرم. تشکر مخصوص من به لنا عباسی است که هر جا ناامید شده بودم مرا امید داد و هر جا مایوس شده بودم به پیش راند و در ضمن پای ثابت و همیشگی نظریه ها و خیال پردازی های من بوده است. از بابت وجود او خوش حال ام!

در نهایت و مهم تر از همه می بایست از دلشاد، نوید، کیان و مادر بزرگام تشکر کنم که همیشه پشتیبان من بودند و به ترین شرایط ممکن را برای فعالیت علمی برای ام فراهم کردند. متأسفانه نوشتار قادر به بیان احساس من نسبت به آن ها نیست و در نتیجه بیش از این سعی در بیان اش نمی کنم.

متشکرم!

فهرست

۱	مقدمه و پیشینه‌ی پژوهش.....
۱-۱	مقدمه
۳	۲-۱ معماری‌های رفتارگرا.....
۷	۳-۱ کارهای مرتبط.....
۷	۱-۳-۱ طراحی خودکار عامل‌های رفتارگرا.....
۱۰	۲-۳-۱ یادگیری تقویتی سلسله مراتبی.....
۱۶	۲ معرفی مساله.....
۲۶	۳ یادگیری ساختار.....
۲۶	۱-۳ مقدمه
۲۸	۲-۳ یادگیری ساختار با بازنمایی مرتبه صفر.....
۲۸	۱-۲-۳ بازنمایی.....
۳۲	۲-۲-۳ تخصیص اعتبار و به روز رسانی بازنمایی.....
۳۵	۳-۳ یادگیری ساختار با بازنمایی مرتبه یک.....
۳۵	۱-۳-۳ بازنمایی.....
۳۸	۲-۳-۳ تخصیص اعتبار و به روز رسانی بازنمایی.....
۴۱	۴ یادگیری رفتار.....
۴۱	۱-۴ مقدمه
۴۲	۲-۴ تجزیه ارزش به اجزای رفتاری.....
۴۴	۳-۴ تخصیص اعتبار به رفتارها.....
۴۴	۱-۳-۴ سیستم چندعامله به عنوان یک جبر بول.....
۴۷	۲-۳-۴ تخصیص اعتبار به رفتارها در معماری Subsumption.....
۵۲	۴-۴ بهینه‌گی، مجموعه عمل‌ها و به روز رسانی ارزش.....
۵۹	۵-۴ یادگیری هم‌زمان رفتار و ساختار.....
۶۱	۵ تکامل رفتار.....
۶۱	۱-۵ مقدمه
۶۶	۲-۵ تکامل رفتارها.....
۶۸	۱-۲-۵ تکامل هم‌کارانه‌ی رفتارها.....
۶۹	۲-۲-۵ اشتراک سازگاری یک‌نواخت.....

۱۲	اپراتورهای ژنتیکی (۴-۲-۵)
۷۳	الگوریتم ممیک (۳-۵)
۸۲	آزمایش‌ها (۶)
۸۲	(۱-۶) مساله مجرد
۸۴	(۱-۱-۶) یادگیری ساختار
۸۶	(۲-۱-۶) یادگیری رفتار و یادگیری همزمان رفتار/ساختار
۹۲	(۳-۱-۶) تکامل رفتار و یادگیری ساختار
۱۰۱	(۲-۶) بلند کردن جسم چند-روباتی
۱۰۳	(۱-۲-۶) یادگیری ساختار
۱۰۹	(۲-۲-۶) یادگیری رفتار و یادگیری همزمان رفتار/ساختار
۱۲۱	(۳-۲-۶) تکامل رفتار و یادگیری ساختار
۱۳۴	تحلیل احتمالاتی معماری PPSSA (۷)
۱۳۴	(۱-۷) مقدمه
۱۳۴	(۲-۷) احتمال کنترل کننده بودن رفتارها
۱۴۷	(۳-۷) تحلیل‌ای بر سرعت یادگیری در معماری PPSSA
۱۴۷	(۱-۳-۷) توزیع دو جمله‌ای منفی
۱۵۸	تاثیر عدم قطعیت سیگنال تقویت بر تابع ارزش و سیاست عامل (۸)
۱۵۸	(۱-۸) مقدمه
۱۶۰	(۲-۸) تاثیر عدم قطعیت سیگنال تقویت بر توابع ارزش
۱۶۴	(۳-۸) تاثیر عدم قطعیت سیگنال تقویت بر سیاست عامل
۱۶۷	(۴-۸) آزمایش‌ها
۱۷۰	(۵-۸) نتیجه‌گیری
۱۷۱	نتیجه‌گیری و پیش‌نهاد (۹)
۱۷۵	(۱۰) مراجع
۱۸۱	ضمیمه الف) واژه‌نامه فارسی-انگلیسی
۱۸۴	ضمیمه ب) واژه‌نامه انگلیسی-فارسی

فهرست جداول

- جدول ۴-۱. الگوریتم کلی برای نمایش عمل کرد عامل ای که ساختار و رفتارش را یاد می گیرد.
- جدول ۵-۱. الگوریتم ممینک. رفتارها تکامل می یابند، ساختار یاد گرفته می شود و ممها انتقال دانش ساختار مناسب را بر عهده دارند.
- جدول ۶-۱. شرایط آزمایش برای آزمون های یادگیری
- جدول ۶-۲. شرایط آزمایش برای آزمون های تکامل / الگوریتم ممینک

فهرست شکل ها

- شکل ۱-۱. تجزیه کارکردگرایانه ی یک وظیفه (روی کرد هوش مصنوعی کلاسیک)
- شکل ۱-۲. تجزیه رفتارگرایانه ی یک وظیفه (روی کرد طراحی رفتارگرا)
- شکل ۲-۱. نمونه ای از معماری Subsumption موازی
- شکل ۴-۱. زیرفضاهای تحریک پذیر رفتارهای B_1 و B_2 در فضای S و تناظر آنها با S'_1 و S'_2 . توجه کنید که زیرفضاهای تحریک پذیر ممکن است هم پوشانی داشته باشند.
- شکل ۵-۱. ساخت عامل جدید با استفاده از اطلاعات ژنتیکی و اطلاعات فرهنگی
- شکل ۶-۱. (مساله ی مجرد) چهار رفتار مختلف مساله ی مجرد. به مناطق مختلف تحریک و درستی یا نادرستی پاسخ در آن مناطق دقت کنید.
- شکل ۶-۲. (مساله ی مجرد) مقایسه بین روش های یادگیری ساختار (مبتنی بر بازنمایی مرتبه صفر و بازنمایی مرتبه یک)
- شکل ۶-۳. (مساله ی مجرد) فضای مساله برای آزمون های یادگیری رفتار / ساختار
- شکل ۶-۴. (مساله ی مجرد) مقایسه پاداش دریافتی برای یادگیری رفتار (ساختار ثابت) و یادگیری هم زمان ساختار و رفتار.

شکل ۵-۶ (مساله‌ی مجرد) مقایسه احتمال خطای یادگیری رفتار (ساختار ثابت) و یادگیری هم‌زمان ساختار و رفتار.

شکل ۶-۶ (مساله مجرد) نواحی یادگرفته شده (اپیزود ۱۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

شکل ۶-۷ (مساله مجرد) نواحی یادگرفته شده (اپیزود ۵۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

شکل ۶-۸ (مساله مجرد) نواحی یادگرفته شده (اپیزود ۱۵۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

شکل ۶-۹ (مساله مجرد) نواحی یادگرفته شده (اپیزود ۳۰۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

شکل ۶-۱۰ (مساله مجرد) مقایسه‌ی بین متوسط و حداکثر سازگاری جمعیت به ازای روش‌های مختلف اشتراک سازگاری همراه با یادگیری ساختار. اشتراک سازگاری یک‌نواخت (آبی)، اشتراک سازگاری ارزش‌محور (سیاه). خط سبز متوسط سازگاری معماری‌ای با ساختار و رفتار کاملاً تصادفی است (با احتمال انتخاب عمل و NA برابر) و خط قرمز حداکثر سازگاری دست‌یافتنی است. خطوط پررنگ متوسط سازگاری جمعیت را نشان می‌دهند و خطوط نقطه‌چین، حداکثر سازگاری را.

شکل ۶-۱۱ (مساله مجرد) مقایسه‌ی متوسط و حداکثر سازگاری به ازای روش‌های مختلف‌ای که از اشتراک سازگاری یک‌نواخت استفاده می‌کنند: تکامل رفتارها و یادگیری ساختار بدون استفاده از مم (آبی)، تکامل رفتارها و یادگیری ساختار با استفاده از مم (سیاه)، تکامل رفتارها و ساختار ثابت (بنفش). خطوط پررنگ متوسط سازگاری جمعیت و خطوط نقطه‌چین حداکثر سازگاری را نشان می‌دهند. حداکثر سازگاری ممکن با خط قرمز نشان داده شده است.

شکل ۶-۱۲ (مساله‌ی مجرد) مقایسه‌ی چگالی احتمال سازگاری برای تکامل با اشتراک سازگاری یک‌نواخت در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری اولیه ساختار به کمک مم (آبی تیره)، عامل‌هایی با ساختار اولیه تصادفی (سبز) و عامل‌هایی با ساختار ثابت (قرمز) انجام گرفته است.

شکل ۶-۱۳ (مساله‌ی مجرد) مقایسه‌ی توزیع احتمال تجمعی ($P\{Fitness \leq f\}$) سازگاری برای تکامل با اشتراک سازگاری یک‌نواخت در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری اولیه ساختار به کمک مم (سیاه)، عامل‌هایی با ساختار اولیه تصادفی (آبی) و عامل‌هایی با ساختار اولیه ثابت (بنفش) انجام گرفته است. تمایل نمودار به سمت راست نشان‌گر بالاتر بودن احتمال تولید عامل‌های سازگارتر است.

شکل ۱۴-۶. (مساله مجرد) مقایسه‌ی متوسط و حداکثر سازگاری به ازای روش‌های مختلف‌ای که از اشتراک سازگاری‌ی ارزش‌محور استفاده می‌کنند: تکامل رفتارها و یادگیری ساختار بدون استفاده از مم (آبی)، تکامل رفتارها و یادگیری ساختار با استفاده از مم (سیاه)، تکامل رفتارها و ساختار ثابت (بنفش). خطوط پررنگ متوسط سازگاری‌ی جمعیت و خطوط نقطه‌چین حداکثر سازگاری را نشان می‌دهند. حداکثر سازگاری‌ی ممکن با خط قرمز نشان داده شده است.

شکل ۱۵-۶. (مساله‌ی مجرد) مقایسه‌ی چگالی احتمال سازگاری برای تکامل با اشتراک سازگاری ارزش‌محور در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری‌ی اولیه ساختار به کمک مم (آبی تیره)، عامل‌هایی با ساختار اولیه‌ی تصادفی (سبز) و عامل‌هایی با ساختار ثابت (قرمز) انجام گرفته است.

شکل ۱۶-۶. (مساله‌ی مجرد) مقایسه‌ی توزیع احتمال تجمعی ($P\{Fitness \leq f\}$) سازگاری برای تکامل با اشتراک سازگاری ارزش‌محور در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری‌ی اولیه ساختار به کمک مم (سیاه)، عامل‌هایی با ساختار اولیه تصادفی (آبی) و عامل‌هایی با ساختار اولیه ثابت (بنفش) انجام گرفته است. تمایل نمودار به سمت راست نشان‌گر بالاتر بودن احتمال تولید عامل‌های سازگارتر است.

شکل ۱۷-۶. سه روبات در حال بلند کردن جسم‌ای بزرگ

شکل ۱۸-۶. (مساله‌ی بلندکردن جسم) مقایسه متوسط سیگنال تقویتی دریافتی بین روش‌های یادگیری ساختار مبتنی بر بازنمایی مرتبه صفر (ZO) و بازنمایی مرتبه یک (FO)، ساختار از پیش طراحی شده و ساختار اتفاقی. رفتارها توسط انسان طراحی شده‌اند.

شکل ۱۹-۶. (مساله‌ی بلندکردن جسم) مقایسه معیار رفتارگرایانه بین روش‌های یادگیری ساختار مبتنی بر بازنمایی مرتبه صفر (ZO) و بازنمایی مرتبه یک (FO)، ساختار از پیش طراحی شده و ساختار اتفاقی. رفتارها توسط انسان طراحی شده‌اند.

شکل ۲۰-۶. (مساله‌ی بلندکردن جسم) مقایسه متوسط سیگنال تقویت دریافتی بین یادگیری ساختار (رفتار طراحی‌شده)، یادگیری رفتار (ساختار طراحی‌شده) و یادگیری هم‌زمان ساختار و رفتار.

شکل ۲۱-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی سیگنال تقویت دریافتی

($P\{\text{average gained reward} \leq x\}$) در فاز یادگیری (بین اپیزود ۱ تا ۵۰. کاوش برقرار است) به

ازای ترکیب‌های مختلف روش‌های یادگیری.

شکل ۲۲-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی سیگنال تقویت دریافتی

($P\{\text{average gained reward} \leq x\}$) در فاز آزمون (بین اپیزود ۵۱ تا ۱۰۰. کاوش پایان یافته است)

به ازای ترکیب‌های مختلف روش‌های یادگیری.

شکل ۲۳-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی معیار رفتارگرایانه

($P\{\text{behavioral performance} \leq x\}$) در فاز یادگیری (بین اپیزود ۱ تا ۵۰. کاوش برقرار است) به

ازای ترکیب‌های مختلف روش‌های یادگیری.

شکل ۲۴-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی معیار رفتارگرایانه

($P\{\text{behavioral performance} \leq x\}$) در فاز آزمون (بین اپیزود ۵۱ تا ۱۰۰. کاوش پایان یافته

است) به ازای ترکیب‌های مختلف روش‌های یادگیری.

شکل ۲۵-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی پاداش دریافت شده در فاز یادگیری (اپیزود ۱ تا ۵۰.

کاوش برقرار است) به ازای روش‌های مختلف یادگیری.

شکل ۲۶-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی پاداش دریافت شده در فاز آزمون (اپیزود ۵۱ تا ۱۰۰.

کاوش پایان یافته است) به ازای روش‌های مختلف یادگیری.

شکل ۲۷-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی معیار رفتارگرایانه در فاز یادگیری (اپیزود ۱ تا ۵۰.

کاوش برقرار است) به ازای روش‌های مختلف یادگیری.

شکل ۲۸-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی معیار رفتارگرایانه در فاز آزمون (اپیزود ۵۱ تا ۱۰۰.

کاوش پایان یافته است) به ازای روش‌های مختلف یادگیری.

شکل ۲۹-۶. (مساله‌ی بلندکردن جسم) نمونه‌ای از مسیر حرکت نقاط تماس روبات-جسم، زاویه‌ی کج‌شده‌گی به

هنگام بلند کردن جسم، و رفتارهای کنترل‌کننده‌ی روبات در هر واحد زمانی. این مسیر حاصل

چندین اپیزود یادگیری هم‌زمان رفتار و ساختار است. تناظر رفتارها و شماره‌های روی نمودار پایین

بدین‌گونه است:

0 (No Behavior), 1 (Push More), 2 (Don't Go Fast), 3 (Stop), 4 (Hurry up), 5

(Slow down).

شکل ۳۰-۶. (مساله‌ی بلند کردن جسم) متوسط سازگاری پنج اپیزود نهایی به ازای طراحی‌های مختلف: (۱) تکامل رفتارها (اشتراک سازگاری ی یک‌نواخت) و یادگیری ساختار (آبی)، (۲) تکامل رفتارها (اشتراک سازگاری ارزش‌محور) و یادگیری ساختار (سیاه)، (۳) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۴) رفتارها و ساختار طراحی‌شده (قرمز). خطوط نقطه‌چین اطراف حالت‌های طراحی‌شده (۳ و ۴) ناحیه‌ی یک انحراف معیار از متوسط را نشان می‌دهند.

شکل ۳۱-۶. (مساله‌ی بلند کردن جسم) متوسط سازگاری پنج اپیزود نهایی و متوسط سازگاری کل دوره‌ی زندگی عامل در طراحی‌های مختلفی که از اشتراک سازگاری یک‌نواخت استفاده می‌کنند: (۱) تکامل رفتارها و یادگیری ساختار بدون کمک مم (آبی)، (۲) تکامل رفتارها و یادگیری ساختار با کمک مم (سیاه)، (۳) تکامل رفتارها و ساختار طراحی‌شده (بنفش)، (۴) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۵) رفتارها و ساختار طراحی‌شده (قرمز). خطوط پررنگ بیان‌گر متوسط سازگاری در پنج اپیزود انتهایی زندگی عامل هستند و خطوط نقطه‌چین سازگاری کل زندگی عامل را نشان می‌دهند. با وجود آن‌که عمل‌کرد آخرین اپیزودهای همه‌ی عامل‌ها تقریباً یک‌سان است اما عمل‌کرد کل دوره‌ی زندگی عامل‌هایی که از اطلاعات مم استفاده می‌کنند بسیار بالاتر است.

شکل ۳۲-۶. (مساله‌ی بلند کردن جسم) مقایسه‌ی چگالی احتمال سازگاری عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ی یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.

شکل ۳۳-۶. (مساله‌ی بلند کردن جسم) مقایسه‌ی چگالی احتمال سازگاری کل دوره‌ی زندگی عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ی یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.

شکل ۳۴-۶.

(مساله‌ی بلندکردن جسم) مقایسه‌ی احتمال تجمعی سازگاری ($P\{Fitness \leq f\}$) عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ی یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (سیاه)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (آبی) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (بنفش) انجام شده است. خطوط نقطه‌چین سازگاری کل دوره‌ی زندگی عامل را نشان می‌دهند. هر چه توزیع بیش‌تر به سمت راست گرایش داشته باشد، شانس ایجاد عامل‌های خیلی خوب بیش‌تر است.

شکل ۳۵-۶.

(مساله‌ی بلند کردن جسم) متوسط سازگاری پنج ایزود نهایی و متوسط سازگاری کل دوره‌ی زندگی عامل در طراحی‌های مختلفی که از اشتراک سازگاری ارزش‌محور استفاده می‌کنند: (۱) تکامل رفتارها و یادگیری ساختار بدون کمک مم (آبی)، (۲) تکامل رفتارها و یادگیری ساختار با کمک مم (سیاه)، (۳) تکامل رفتارها و ساختار طراحی‌شده (بنفش)، (۴) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۵) رفتارها و ساختار طراحی‌شده (قرمز). خطوط پررنگ بیان‌گر متوسط سازگاری در پنج ایزود انتهای زندگی عامل هستند و خطوط نقطه‌چین سازگاری کل زندگی عامل را نشان می‌دهند. با وجود آن‌که عمل‌کرد آخرین ایزودهای همه‌ی عامل‌ها تقریباً یک‌سان است اما عمل‌کرد کل دوره‌ی زندگی عامل‌هایی که از اطلاعات مم استفاده می‌کنند بالاتر است.

شکل ۳۶-۶.

(مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.

شکل ۳۷-۶.

(مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری کل دوره‌ی زندگی عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.

شکل ۳۸-۶.

(مساله‌ی بلندکردن جسم) مقایسه‌ی احتمال تجمعی سازگاری ($P\{Fitness \leq f\}$) عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از

اطلاعات مهم استفاده کرده‌اند (سیاه)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (آبی) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (بنفش) انجام شده است. خطوط نقطه‌چین سازگاری کل دوره‌ی زندگی عامل را نشان می‌دهند. هر چه توزیع بیش‌تر به سمت راست گرایش داشته باشد، شانس ایجاد عامل‌های خیلی خوب بیش‌تر است.

شکل ۷-۱. تاثیر بعد خروجی بر احتمال کنترل‌کننده بودن. هر چه بعد خروجی بزرگ‌تر باشد، این احتمال کم‌تر می‌شود.

شکل ۷-۲. تاثیر تعداد رفتارها بر احتمال کنترل‌کننده بودن. هر چه تعداد رفتارها بیش‌تر باشد، رفتارهای پایین دارای احتمال کنترل‌کننده بودن کم‌تری هستند.

شکل ۷-۳. تاثیر η بر احتمال کنترل‌کننده بودن. در مقادیر نزدیک به ۱، η ، با کاهش مقدار نمودارها تخت‌تر می‌شوند.

شکل ۷-۴. تاثیر η بر احتمال کنترل‌کننده بودن. مقادیر کوچک η باعث افزایش شدید احتمال کنترل‌کننده بودن لایه‌های پایین و کم شدن آن احتمال برای لایه‌های بالا می‌شود. با کاهش مقدار نمودارها تخت‌تر می‌شوند.

شکل ۷-۵. تاثیر بعد خروجی بر احتمال کنترل‌کننده بودن در مدل کاهش حسابی. هر چقدر بعد خروجی بزرگ‌تر باشد، این احتمال کم‌تر می‌شود.

شکل ۷-۶. تاثیر تعداد رفتارها بر احتمال کنترل‌کننده بودن. هر چه تعداد رفتارها بیش‌تر باشد، رفتارهای پایین دارای احتمال کنترل‌کننده بودن کم‌تری هستند.

شکل ۷-۷. تاثیر ε بر احتمال کنترل‌کننده بودن. با افزایش ε ، منحنی‌ها تخت‌تر می‌شوند.

شکل ۷-۸. چگالی احتمال تعداد تلاش لازم برای ۵ بار کنترل‌کننده بودن رفتاری با احتمال کنترل‌کننده بودن ۰.۱۵، در هر واحد زمانی.

شکل ۷-۹. احتمال بیش از ۵ بار کنترل‌کننده بودن برای معماری دو لایه

- شکل ۷-۱۰. چگالی احتمال بیش از ۵ بار کنترل کننده بودن برای معماری دو لایه
- شکل ۷-۱۱. احتمال بیش از ۵ بار کنترل کننده بودن برای معماری سه لایه
- شکل ۷-۱۲. چگالی احتمال بیش از ۵ بار کنترل کننده بودن برای معماری دو لایه
- شکل ۸-۱. نمونه‌ای از مساله‌ی MDP که کران‌های ارزیابی شده دقیق‌اند.
- شکل ۸-۲. خطای $\| \Delta Q(s, a) \|_{\infty}$ به ازای مقادیر γ مختلف
- شکل ۸-۳. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.1$
- شکل ۸-۴. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.5$
- شکل ۸-۵. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.9$
- شکل ۸-۶. کران بالا و پایین نسبت احتمالات عامل با سیگنال تقویت نادقیق به احتمالات عامل با سیگنال تقویت اصلی به ازای مقادیر مختلف γ (آبی: $\gamma=0.1$ ؛ مشکی: $\gamma=0.5$ ؛ قرمز: $\gamma=0.9$).
- شکل ۸-۷. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.1$
- شکل ۸-۸. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.5$
- شکل ۸-۹. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.9$

(۱) مقدمه و پیشینه‌ی پژوهش

(۱-۱) مقدمه

عمومیت استفاده از عامل‌های^۱ خودکاری که در محیط‌هایی مانند کارخانه یا منازل به کار می‌روند روز به روز بیش‌تر می‌شود. به تدریج روبات‌ها از آزمایش‌گاه‌ها خارج شده و وارد محیط‌های واقعی می‌شوند. این فرآیند نیازمند وجود مکانیزم‌های تصمیم‌گیری مقاوم‌ای است که بتواند به خوبی از عهده‌ی نایقینی‌ها و خطاهای پیش‌آمده برای روبات برآید. روبات‌ای که در محیط واقعی قرار می‌گیرد با بسیاری از عوامل اختلال‌زا مانند نویز سنسور، خطای محرک^۲ و پیچیدگی و تغییر دایم محیط مواجه است. هم‌چنین روبات در اکثر موارد به همه‌ی اطلاعات لازم برای تصمیم‌گیری دسترسی ندارد و میزان اطلاعات‌ای که سنسورها به او می‌دهند محدود است. تصمیم‌گیری در چنین شرایطی بسیار دشوار است و تلاش‌های اولیه‌ای که خواهان به‌کارگیری روش‌های کلاسیک هوش مصنوعی برای حل این‌گونه مسایل بودند با شکست مواجه شده‌اند. خوش‌بختانه معماری‌های رفتارگرا^۳ که از اواسط دهه‌ی هشتاد میلادی مطرح شده‌اند بسیاری از این مشکلات عملی را حل کرده‌اند. معماری‌های رفتارگرا با تجزیه‌ی کل مساله به تعدادی رفتار مجزا (به جای تقسیم به تعدادی کارکرد^۴) بسیاری از مسایل دشوار پیشین را حل می‌کنند. با این‌همه هم‌چنان طراحی معماری مناسب‌ای برای حل مساله‌ای خاص چندان ساده نیست و در عمل طراح با سعی و خطای بسیار به تدریج به حل مطلوب نزدیک خواهد شد. این پرسش مطرح می‌شود که آیا می‌توان بخش‌ای از این فعالیت وقت‌گیر را به خود عامل واگذار کرد تا با توجه به راهنمایی‌های طراح بتواند به معماری مطلوب برسد؟ تلاش برای پاسخ مثبت دادن به چنین پرسشی انگیزه‌ی اصلی‌ی این پژوهش بوده است. هدف اصلی‌ی این پژوهش ایجاد مکانیزم‌هایی برای طراحی خودکار معماری‌های رفتارگرا است. به عبارت دیگر، قصد داشتیم تا روش‌های کلی‌ای ایجاد کنم تا به وسیله‌ی آن طراح بتواند بخش عمده‌ای از دشواری فرآیند طراحی را به خود عامل واگذارد. در این

¹ Agents

² Actuator

³ Behavior-based Architectures

⁴ Function

پژوهش ایده‌های گوناگونی را برای خودکارسازی طراحی معماری‌های رفتارگرایی سلسله مراتبی پیش‌نهاد می‌کنم.

برای خودکار کردن فرآیند طراحی لازم است معیاری برای برتری طراحی بر طرح دیگر وجود داشته باشد. این معیار که توسط طراح مشخص می‌شود بیان‌گر دیدگاه او به راه حل مناسب مساله است. به عنوان مثال در اکثر موبایل روبات‌ها برخورد کردن با اجسام رفتار مطلوب‌ای نیست. طراح می‌تواند این خواست را به صورت امتیاز منفی‌ای مشخص کند: هر برخورد معادل یک واحد امتیاز منفی باشد. آن‌گاه با این معیار، طراحی به‌تر است که امتیاز منفی کم‌تری بگیرد و در نتیجه به‌ترین طرح ممکن معادل کنترل‌ری است که کم‌ترین برخورد را موجب شود. با این دیدگاه، فرآیند طراحی معادل با بهینه‌سازی معیار مشخص شده توسط طراح است.

در این پژوهش برای حل مساله‌ی بهینه‌سازی یاد شده دو رویکرد متفاوت را بررسی کرده‌ام. اولین رویکرد، استفاده از یادگیری تقویتی^۵ بوده است. روبات با استفاده از سیگنال تقویت^۶ - که به صورت پاداش یا تنبیه‌ای برای عامل خواهد بود- سعی در ایجاد به‌ترین معماری رفتارگرایی ممکن می‌کند. این فرآیند - که در آینده به طور دقیق به آن خواهیم پرداخت- در دو سطح تطبیق ساختار معماری (فصل ۳) و تطبیق رفتارهای قرارگرفته در معماری (فصل ۴) صورت می‌پذیرد. رویکرد دیگری که برای حل مساله‌ی بهینه‌سازی مطرح شده، بهره‌جویی از تکامل مصنوعی^۷ است. با استفاده از تکامل، جمعیت‌ای از عامل‌ها سعی در پیدا کردن به‌ترین معماری ممکن می‌کنند. در پژوهش حاضر فرآیند تکامل در سطح تطبیق رفتارها و به صورت تکامل هم‌کارانه انجام می‌شود (فصل ۵). علاوه بر این دو نگاه، استفاده از ایده‌های فرهنگی و انتقال دانش از طریق فرهنگ مشترک بین عامل‌ها با معرفی مفهوم‌ای به نام مم^۸ و الگوریتم‌های ممیک^۹ نیز مورد مطالعه قرار گرفته است (فصل ۵). نتایج این دو رویکرد را در فصل مربوط به آزمایش‌ها (فصل ۶) آورده‌ام. در فصل ۷ به تحلیل احتمالاتی معماری

^۵ Reinforcement Learning

^۶ Reinforcement Signal

^۷ Artificial Evolution

^۸ Meme

^۹ Memetic Algorithm



شکل ۱-۱. تجزیه کارکردگرایانه‌ی یک وظیفه (روی کرد هوش مصنوعی کلاسیک)

سلسله‌مراتبی به کار رفته در این پایان‌نامه می‌پردازم و خواص آن را مطالعه می‌کنم. در فصل ۸ نیز به تاثیر وجود عدم قطعیت در سیگنال تقویت بر تابع ارزش و تابع سیاست عامل می‌پردازم. در بخش‌های پیش رو بیش‌تر درباره‌ی معماری‌های رفتارگرا و هم‌چنین پژوهش‌های پیشین‌ای که برای خودکار کردن طراحی معماری‌های رفتارگرا صورت گرفته است خواهم نوشت.

۲-۱) معماری‌های رفتارگرا

برخلاف عامل‌هایی که در محیط‌های مجرد^{۱۰} تصمیم‌گیری می‌کنند (مانند عامل جستجوگر ماز، عامل شطرنج‌باز و ...)، عامل جسم‌دار^{۱۱} در محیط واقعی قرار گرفته^{۱۲} با بسیاری از عدم قطعیت‌ها روبرو است: ورودی اطلاعات چنان عامل‌ای نویزی و نادقیق است، در اکثر مواقع دانش‌اش نسبت به محیط بسیار جزئی و ناکامل است (از دید عامل، محیط به طور جزئی رویت‌پذیر^{۱۳} است)، عمل‌ای که انتخاب می‌کند الزاما به همان صورت اجرا نمی‌شود و هم‌چنین معمولا محیط پیچیده و متغیر با زمان است. همه‌ی این‌ها باعث می‌شود که روی‌کرد متداول هوش مصنوعی کلاسیک که مبتنی بر ادراک^{۱۴} محیط، مدل مجردسازی از محیط، برنامه‌ریزی^{۱۵} در آن مدل مجرد و سپس اجرای فرمان‌ها باشد (شکل ۱-۱) غیرعملی باشد. عدم قطعیت محیط از دید عامل آن چنان دقت مدل ساخته شده را پایین می‌آورد که برنامه‌ریزی انجام گرفته تقریبا بی‌فایده خواهد شد. طراحی رفتارگرا بسیاری از این مشکلات را حل

¹⁰ Abstract

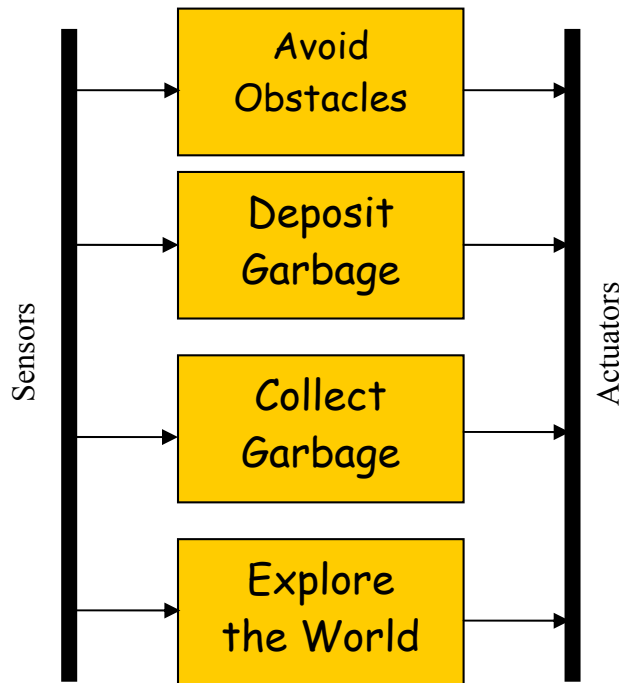
¹¹ Embodied

¹² Situated

¹³ Partial Observable

¹⁴ Perception

¹⁵ Planning



شکل ۲-۱. تجزیه رفتارگرایانه‌ی یک وظیفه (روی کرد طراحی رفتارگرا)

می‌کند و به کمک آن دیدگاه روبات‌هایی طراحی شده‌اند که تا پیش از آن ممکن نبودند ([Brooks86]، [Brooks91a]، [Brooks91b]، [Brooks91c] و [Arkin98]). برای نمونه‌ای از کارهای موفق‌ای که با استفاده از معماری‌ای رفتارگرا انجام گرفته است به [Brooks89]، [Wang96]، [Mataric98]، [Parker98]، [Nili01] مراجعه کنید.

این که طراحی رفتارگرا شامل چه مولفه‌هایی است و چه زمینه‌ی فلسفه‌ای دارد به طور دقیق مورد توافق همه‌ی پژوهش‌گران نیست. با این وجود شاید بتوان گفت که مهم‌ترین ویژگی آن، نوع تجزیه‌ی رفتار کلی‌ی عامل به تعدادی رفتار جزئی‌تر است. این رفتارهای جزئی هستند که در نهایت رفتار کلی‌ی عامل را ایجاد می‌کند. رفتارهای جزئی با تعامل با محیط و همچنین تعامل با یک‌دیگر باعث تظاهر^{۱۶} رفتار هوش‌مندانه می‌شوند. در این دیدگاه به جای نگاه‌ای کارکردگرایانه^{۱۷} به اجزای تصمیم‌گیری (کارکرد درک محیط، کارکرد مدل‌سازی، کارکرد برنامه‌ریزی و ...)، نگاه‌ای رفتارگرایانه^{۱۸} به

¹⁶ Emergence

¹⁷ Functionalism

¹⁸ Behaviorism

مساله داریم. نگاه‌ای که با دیدی عینی^{۱۹} نه به ساز و کار داخلی ذهن که به نتیجه‌ی تصمیم ذهنی‌ی عامل در محیط -با توجه به همه‌ی قیده‌های فیزیکی‌ی عامل و محیط- می‌پردازد.

برای روشن‌تر شدن قضیه، فرض کنید که قصد ساخت روبات‌ای را داریم که آشغال‌ها را از روی زمین جمع می‌کند و به مکان مشخص‌ای می‌برد. هم‌چنین می‌دانیم که روبات نباید به اجسام محیط برخورد کند. دیدگاه کلاسیک (و کارکردگرایانه) ابتدا سعی در درک محیط با استفاده از سنسورهای روبات می‌کند و سپس با استفاده از اطلاعات استخراج شده به ساختن مدل‌ای مجرد از دنیای اطرافش تلاش می‌کند. چنین کاری ساده نیست چون عدم قطعیت سنسورها باعث پیچیده و یا حتی غیرممکن شدن عمل تشخیص دقیق ساختار محیط و مکان آشغال‌های دور و بر می‌شوند. پس از انجام این کار -در صورت امکان- می‌بایست در فضای بسیار بزرگ ایجاد شده به گونه‌ای مسیر روبات را برنامه‌ریزی کرد تا نه تنها همه‌ی آشغال‌ها جمع شوند بلکه قیده‌های دیگری چون برخورد نکردن با موانع نیز برقرار شود. به دلیل ابعاد بزرگ فضای برنامه‌ریزی، چنین کاری الزاما خیلی ساده نیست. مهم‌تر از آن به دلیل مجرد بودن ورودی‌ی برنامه‌ریز، در صورت‌ای که مرحله‌ی مدل‌سازی دقیق انجام نشده باشد -که در اکثر موارد این‌چنین است- خطای مدل‌سازی باعث برنامه‌ریزی در محیطی می‌شود که متناظری در دنیای واقعی ندارد. حتی در صورت‌ای که بتوان از این موانع گذشت، نتیجه‌ی نهایی نه تنها دقت بالا و مورد انتظار روش‌های کلاسیک را ندارد بلکه به دلیل کند بودن عملیات انجام شده در تجزیه‌ی کارکردگرایانه‌ی وظیفه، عامل بسیار کند و لخت خواهد بود. چنین عمل‌کردی برای مسایل دنیای واقعی قابل قبول نیست.

در دیدگاه طراحی رفتارگرا، وظیفه‌ی مورد نظربه تعدادی جزء رفتاری تجزیه می‌شود (شکل ۱-۲). جزء رفتاری آن چیزی است که ناظر بیرونی به هنگام مشاهده‌ی عامل‌ای که آن مساله را به درستی حل می‌کند آن را یک رفتار مجزا تشخیص می‌دهد. مثلا اگر عامل‌ای -مثلا انسان- عمل مورد نظر را به درستی انجام دهد، ناظر رفتارهای "به سمت آشغال حرکت کردن"، "دوری جستن از اجسام"، "دور زدن موانع"، "به سمت هدف حرکت کردن" و ... را در فعالیت آن عامل تشخیص می‌دهد. توجه کنید که این‌ها الزاما آن چیزی نیستند که در ذهن عامل می‌گذرد اما از آن‌جا که ما به ذهن عامل -مثلا

¹⁹ Objective

انسان - دست‌رسی‌ی مستقیم نداریم، تنها فرآیند عینی و اندازه‌پذیرمان - در تعریف غیرریاضی‌ی آن - همین رفتارهای مشخص شده است. حال اگر بیابیم و هر کدام از این رفتارها را به طور مجزا تولید کنیم، مجموعه‌ای از رفتارهایی خواهیم داشت که با استفاده از آن رفتارها و انتخاب صحیح و به موقع‌شان می‌توان رفتار کلی‌ی عامل را بازتولید نمود. این کار - یعنی تولید رفتار - فرآیندی است به ظاهر متفاوت با فلسفه‌ی رفتارگرایی. بعضی از پژوهش‌گران (مثلاً [Matarić01]) اعتقاد دارند که طراحی رفتارگرا نمونه‌ای از رفتارگرایی نیست. دلیل آن نظر این است که فلسفه‌ی رفتارگرایی اجازه‌ی گمانه‌زنی درباره‌ی ماهیت داخلی ذهن منجر به این رفتار را به تحلیل‌گر رفتار نمی‌دهد. اما شاید توجه به این نکته که هدف رفتارگرایی تحلیل - و نه طراحی - بوده است تا حدی راه‌گشا باشد. طراح معماری‌ی رفتارگرا با الهام از این ایده که حل مساله را به صورت تعدادی رفتار مجزا ببیند و ام‌دار رفتارگرایی است اما به هنگام طراحی‌ی هر کدام از رفتارها از رفتارگرایی جدا شده و به صورت زمینه‌ی پژوهشی‌ی متفاوت‌ای در می‌آید. در نتیجه - به اعتقاد من - طراحی‌ی رفتارگرا با وام‌گیری از بخش‌ای از ایده‌ی رفتارگرایی سعی در حل مساله‌ی ساخت عامل هوشمند می‌کند.

ویژگی‌ی خاص و اساسی‌ی این تجزیه، ساده‌سازی‌ی فرآیند طراحی است. معمولاً در اکثر مسائل، طراحی‌ی یک جزء رفتاری بسیار ساده‌تر از طراحی‌ی سیستم کلاسیک‌ای است که برای حل مساله، همه‌ی ابعاد آن را با هم در نظر می‌گیرد. این کار در مسائل دنیای واقعی بسیار پیچیده و غیرقابل پی‌گیری^{۲۰} است. اما برخلاف آن طراحی‌ی یک رفتار ساده - مانند "دوری جستن از اجسام" - نیازی به تحلیل دقیق و کامل دنیای اطراف ندارد و به جای کارکردن در کل فضای مساله، تنها بخش کوچک ولی مربوطی مانند فاصله با اجسام نزدیک مورد توجه قرار می‌گیرد. پس از طراحی‌ی اجزای رفتاری مورد نیاز می‌بایست شرایط فعال شدن هر کدام از آن رفتارها را مشخص کرد. این که چه زمان‌ای رفتار خاص‌ای می‌بایست فعال شود و کنترل عامل را در اختیار بگیرد، مساله‌ای است که بخش مهم‌ای از فرآیند طراحی رفتارگرا را شامل می‌شود. در اکثر موارد این تصمیم‌گیری به عهده‌ی طراح است و او با روش‌ای پایین به بالا^{۲۱} سعی در پیدا کردن بهترین ترکیب رفتارها می‌کند: طراح ابتدا با انتخاب یک رفتار

²⁰ Intractable

²¹ Bottom up

شروع می‌کند و رفتار کلی‌ی عامل را مشاهده می‌کند. در صورت مطلوب نبودن نتیجه یا شرایط کنترل‌کننده بودن رفتار را تغییر می‌دهد یا رفتار جدیدی ایجاد می‌کند (یا از جعبه ابزار رفتارهای‌اش، رفتار جدیدی را انتخاب می‌کند) تا در مواردی که مجموعه رفتارهای پیشین به خوبی عمل نکرده‌اند وارد عمل شده و باعث رفتار مطلوب عامل شود. این عمل آن‌قدر تکرار می‌شود تا رفتار کلی‌ی مورد نظر به دست آید.

روی‌کرد طراحی‌ی رفتارگرا معمولاً برای طراحی‌ی کنترل‌کننده‌های هوش‌مندی که قرار است برای عامل‌های محیط واقعی عمل کند موفق‌تر از روی‌کرد کلاسیک نتیجه می‌دهد. طراحی‌ی رفتارگرا نه تنها ساده‌تر است بلکه از آن‌جا که فرآیند تصمیم‌گیری‌ی درونی‌ی هر رفتار در اکثر مواقع ساده‌تر و کم محاسبه‌تر از فرآیند لازم برای حلِ کلِ مساله است، کنترلر حاصل بسیار سریع‌تر واکنش می‌دهد. این نکته مخصوصاً با در نظر گرفتن این‌که طراحی‌ی رفتارگرا ذاتاً گسترده^{۲۲} است (هر رفتار یک پروسه‌ی مجزا است) آشکارتر می‌شود.

هدف پژوهش فعلی همان‌طور که پیش‌تر گفته شد خودکار کردن فرآیند طراحی‌ی معماری‌ی رفتارگرا است. سعی کرده‌ام با استفاده از یادگیری و تکامل روش‌ای برای طراحی‌ی خودکار معماری‌ی رفتارگرا – که تا پیش از این توسط طراح طرح می‌شد – معرفی کنم.

۱-۳ کارهای مرتبط

در این بخش مروری بر کارهای مشابه که به نحوی به پژوهش فعلی مربوط است خواهم داشت. در ابتدا مروری بر روش‌های طراحی‌ی خودکار عامل‌های رفتارگرا می‌کنم و سپس بعضی از روش‌های یادگیری‌ی سلسله مراتبی را بررسی خواهم کرد.

۱-۳-۱ طراحی‌ی خودکار عامل‌های رفتارگرا

تقسیم روش‌های ممکن برای طراحی‌ی معماری‌های رفتارگرا به دو دسته‌ی روش‌های مبتنی بر یادگیری و روش‌های مبتنی بر تکامل ویژه‌ی این پژوهش نیست.

²² Distributed

روش‌های مبتنی بر یادگیری، نتایج اولیه‌ی قابل توجه و موفقیت‌آمیزی در طراحی‌ی خودکار سیستم‌های رفتارگرا داشته‌اند. اما اکثر آن‌ها مساله‌ی یادگیری را در فرم چندان کامل و جامع‌ای بررسی نکرده‌اند. هم‌چنین تحلیل ریاضی و دقیق‌ای بر چگونگی‌ی حل مساله صورت نگرفته است و اکثر روش‌ها تنها به بهینه‌سازی قسمت‌ای از معماری پرداخته‌اند. به عنوان نمونه‌ای از کارهای انجام شده می‌توان به [Maes90] اشاره کرد که از یادگیری برای تنظیم پیش‌شرط فعال شدن مجموعه‌ای از رفتارهای از پیش تعیین شده استفاده کرده است. این پژوهش یکی از اولین کارهایی است که در این زمینه صورت گرفته و از روی‌کردی مشابه – اما نه دقیقاً یک‌سان – با آنچه در تئوری‌ی محاسباتی شده‌ی یادگیری تقویتی معمول است بهره جسته (برای شرحی از روش‌های یادگیری تقویتی به [Sutton98] مراجعه کنید). در [Mahadevan92]، از یادگیری برای تغییر رفتارها در معماری‌ای بسیار مشابه با معماری‌ای که در این پژوهش استفاده کرده‌ام – یعنی معماری‌ی Subsumption –ی موازی – استفاده شده است. در آن پژوهش ساختار از پیش ثابت است و هر رفتار دارای سیگنال تقویت مجزایی است. این برخلاف روی‌کرد پژوهش فعلی است که در آن نه معماری ثابت است و نه هر لایه – یا رفتار – دارای سیگنال تقویت مجزایی است. در [Matarić92] از ساختاری کاملاً رفتارگرا برای یادگیری نقشه‌ی توپولوژیک محیط استفاده شده است. در آن کار ساختار معماری ثابت نیست و در صورت لزوم گسترش می‌یابد. افزودن هر رفتار معادل یافتن زمین‌نشان‌ای^{۲۳} جدید است. در [Matarić94] و [Matarić97]، از سیگنال تقویت شکل‌دهی‌شده^{۲۴} و تخمین‌زن پیش‌رفت^{۲۵} برای سرعت بخشیدن به یادگیری در معماری‌ای رفتارگرا بهره گرفته شده است. البته معماری‌ی به‌کار رفته در آن کار سلسله مراتبی نیست و رفتارها در یک سطح با هم رقابت می‌کنند. هم‌چنین در [Michaud98] نیز برای انتخاب رفتارها از روی‌کردی حافظه‌محور^{۲۶} استفاده شده است. بدین معنا که با توجه به تاریخ‌چه‌ی حالت‌هایی که روبات در آن قرار گرفته است و با توجه به نرخ موفقیت روبات در آن حالت‌ها، امتیازی به آن حالت داده می‌شود. این کار شبیه به کارهای مک‌کالوم^{۲۷} در Nearest Sequence Matching

²³ Landmark

²⁴ Shaped Reinforcement Signal

²⁵ Progress Estimator

²⁶ Memory-based

²⁷ McCallum

([McCallum95a]) و Utile Suffix Memory ([McCallum95b]) برای در نظر گرفتن حافظه و تاثیر تاریخچه‌ی حالت‌ها در انتخاب عمل مناسب است. برای مروری بر کارهای مشابه که از یادگیری برای تطبیق در معماری‌های رفتارگرا استفاده کرده باشند به [Matarić01] مراجعه کنید.

روش‌های مبتنی بر تکامل نیز در روباتیک استفاده شده است. در این روش‌ها کنترلر روبات – که معمولاً به صورت یک پارچه^{۲۸} توصیف می‌شوند – با استفاده از میزان سازگاری^{۲۹} عامل استفاده‌کننده از آن کنترلر تکامل می‌یابد. به دلیل یک پارچه بودن اکثر کنترلرهای تکاملی و نبود چندین رفتار مختلف به صورت مشخص و بارز (و نه فقط به شکل نمود خارجی) می‌توان در رفتارگرا بودن این نوع طراحی شک داشت و اکثر کنترلرهای طراحی شده را بیش‌تر جزو خانواده‌ی کنترلرهای راکتیو^{۳۰} در نظر گرفتن. در بسیاری از پژوهش‌ها از کنترلرهای مبتنی بر شبکه‌ی عصبی استفاده می‌شود و پارامترهای تکامل^{۳۱} وزن‌های شبکه‌ی عصبی است. به عنوان نمونه به [Harvey93]، [Cliff93]، [Floreano94]، [Floreano96]، [Urzelai98]، [Floreano99]، [Floreano00]، [Nolfi00]، [Chernova04] مراجعه کنید. به عنوان کارهایی که شباهت بیش‌تری با کار فعلی دارند می‌توان به [Koza94] اشاره کرد. در آن پژوهش از برنامه‌ریزی ژنتیکی^{۳۱} برای تولید معماری‌ی Subsumption مشابه با معماری‌ی فعلی استفاده شده است. هدف استفاده از تکامل^{۳۲} ایجاد کنترلری بوده است که عمل تعقیب دیوار^{۳۲} را به درستی انجام دهد. همچنین [Togelius04] نیز از تکامل لایه‌ای^{۳۳} اضافه‌شونده^{۳۴} برای ایجاد معماری‌ای مشابه با معماری‌ی این پژوهش استفاده کرده است.

با وجود این که روش‌های متداول روباتیک تکاملی به نظر توان‌مند می‌آیند اما ایرادهای مشخص‌ای نیز دارند. به عنوان یکی از بزرگ‌ترین مشکلات آن‌ها می‌توان به سرعت پایین‌شان نسبت به روش‌هایی مبتنی بر یادگیری اشاره کرد. همچنین اکثر کنترلرهای به‌کار رفته در روباتیک تکاملی به

²⁸ Monolithic

²⁹ Fitness

³⁰ Reactive

³¹ Genetic Programming

³² Wall Following

³³ Layered

³⁴ Incremental

دلیل یک پارچه بودنشان قابلیت استفاده‌ی مجدد ندارند. در این باره در فصل ۵ بیش‌تر توضیح خواهیم داد.

۱-۳-۲) یادگیری تقویتی سلسله مراتبی^{۳۵}

ایده‌ی استفاده از یادگیری در ساختارهای سلسله مراتبی ایده‌ی چندان تازه‌ای نیست. در چند سال اخیر پژوهش‌های بسیاری در حوزه‌ای به نام یادگیری تقویتی سلسله‌مراتبی انجام شده است که نتایج موفق‌ای نیز داشته است. از آن‌جا که در این پژوهش از معماری‌ای سلسله مراتبی استفاده شده است، مروری مختصر بر روش‌های متداول یادگیری‌ی آن معماری‌ها مفید می‌نماید.

برای شروع، به این مثال توجه کنید: فرض کنید که موبایل روبات‌ای قصد یادگیری‌ی رفتار خروج از ساختمان را داشته باشد. در روش‌های معمول یادگیری، عمل^{۳۶}‌های ممکن روبات بسیار جزئی هستند: مثلاً حرکت در یکی از چهار جهت اصلی. روبات‌ای که در اتاقی قرار گرفته است برای خروج از آن می‌بایست با حرکت‌های کوچک و پایه‌ای^{۳۷} خود آن‌قدر در محیط به کاوش^{۳۸} بپردازد تا به خروجی‌ی ساختمان برسد. در آن هنگام روبات پاداش‌ای دریافت می‌کند و انتخاب‌های قبلی‌اش تقویت می‌شود. حال اگر طراح این روبات به جای حرکت‌های پایه‌ای، او را به حرکت‌های مرکب‌ای مانند "حرکت به سمت در"، "حرکت در طول راهرو"، "دور زدن از موانع" و ... نیز مجهز کرده باشد، روبات می‌تواند با انتخاب چند حرکت مرکب به خروجی برسد و پاداش لازم را دریافت کند. این حرکت‌های لازم در واقع مجموعه‌ای از حرکت‌های پایه‌ای است که توسط برنامه‌ای ایجاد می‌شود. این مثال بیان‌گر روش متداول‌ای به نام options است [Sutton99]. Options که به آن‌ها Macro Action نیز می‌گویند در اصل تعمیم‌ای بر عمل‌های ساده‌ای است که در یادگیری تقویتی مطرح می‌شود. هر option گسترش زمانی‌ی یک عمل است و به همین دلیل option‌ها را اعمال گسترش-زمانی یافته^{۳۹} نیز می‌نامند. برخلاف انتخاب عمل‌ها در چارچوب MDP^{۴۰} که هر انتخاب تنها باعث انتقال از یک حالت به حالت

³⁵ Hierarchical Reinforcement Learning

³⁶ Action

³⁷ Primitive

³⁸ Exploration

³⁹ Temporally Extended Action

⁴⁰ Markov Decision Problem

دیگر می‌شود، هر option می‌تواند باعث گذشت چند واحد زمانی و انجام چند عمل پایه‌ای^{۴۱} و در نتیجه چندین انتقال حالت متوالی شود. هنگامی که عامل یک option را انتخاب می‌کند، مجموعه‌ای از عمل‌های از پیش تعیین شده پشت سر هم اجرا می‌شود. از دیدگاهی هر option را می‌توان یک زیربرنامه^{۴۲} دانست که مطابق با برنامه‌ای از پیش مشخص به انتخاب اعمال می‌پردازد. هر option حتی می‌تواند درون خود option‌های دیگر را نیز فراخوانی کند و بدین ترتیب تبدیل به روش‌ای بازگشتی^{۴۳} – و در نتیجه سلسله مراتبی – می‌شود.

option‌ها دارای دو ویژگی‌ی مطلوب هستند: اول این که طراح در صورت داشتن دانش از پیش نسبت به رفتارهای مطلوب در آن محیط می‌تواند به طراحی option‌های مناسب آن مساله دست یازد. در این صورت سرعت یادگیری – یا برنامه‌ریزی – بسیار بالا می‌رود. ویژگی دیگر این است که option‌ها با کوچک کردن فضای تصمیم‌گیری باعث ساده‌تر و سریع‌تر شدن برنامه‌ریزی و یادگیری می‌شوند. البته در این بین حداکثر بازده ممکن است کم‌تر شود^{۴۴}.

یادگیری و برنامه‌ریزی با option‌ها – که تئوری‌شان بر پایه‌ی Semi-Markov Decision Process است [Puterman94] – بسیار مشابه با یادگیری در سیستم‌های متداول مبتنی بر MDP است و الگوریتم‌های مشابه با Q-Learning و SARSA نیز در آن وجود دارد.

یکی از مسایل بسیار مهم در این روش سلسله مراتبی پیدا کردن option‌های مناسب است. با وجود آن که در اکثر موارد option‌ها توسط طراح از پیش مشخص می‌شوند، اما چنین option‌هایی الزاما به بالاترین بازده منجر نمی‌شوند: طراح ممکن است درک دقیقی از مساله نداشته باشد و زیربرنامه‌هایی بنویسد که به‌ترین ممکن برای آن مساله نباشد. روش‌های مختلفی برای طراحی خودکار option مطرح شده است. به عنوان مثال در [Precup00]، روش یادگیری intra-option معرفی شده که با استفاده از تابع پاداش ویژه‌ای برای هر زیرهدف^{۴۵} در سیاست^{۴۶} خود option‌ها نیز

⁴¹ Primitive Action

⁴² Subroutine

⁴³ Recursive

⁴⁴ اگر طراح به عامل اجازه دهد تا یادگیری را بر روی مجموعه‌ی حرکات پایه‌اش نیز انجام دهد (و آن‌ها را حذف نکند)، آن‌گاه بازده کم نخواهد شد.

⁴⁵ Subgoal

⁴⁶ Policy

تغییر ایجاد می‌کند. از دیدگاهی بسیار کلی این شبیه به روش یادگیری رفتار این پایان‌نامه است. گرچه ساختار ویژه و سلسله‌مراتبی معماری‌ی این پژوهش تغییرات اساسی‌ای را در روش یادگیری مناسب ایجاد می‌کند. در [McGovern01] نیز روش‌ای برای ایجاد تدریجی زیرهدف‌ها در طول یادگیری معرفی شده است. همچنین در چارچوب Intrinsicly Motivated Reinforcement Learning [Singh04] نیز فرآیندی برای تولید option‌های جدید به هنگام رسیدن به شرایطی که عامل به حالت‌های تازه و ناآشنا می‌رسد پیش‌نهاد شده است.

یکی دیگر از روش‌های یادگیری تقویتی سلسله‌مراتبی، روش Feudal RL است [Dayan93]. در این روش چند لایه، لایه‌های بالاتر – که دید کلی‌تری نسبت به مساله دارند – تصمیم می‌گیرند که کدام لایه‌ی پایین‌تر می‌بایست کنترل عامل را در اختیار بگیرد. لایه‌های بالاتر دیدی کلی‌تر ولی مجردتر نسبت به مساله دارند ولی لایه‌های پایین‌تر – که فرمان‌بردارند – دیدی جزئی ولی دقیق‌تر دارند. روش یادگیری‌ای که در این پژوهش به آن خواهیم پرداخت نیز دارای ساختاری لایه به لایه است اما برخلاف روش Feudal RL، دیدی رفتارگرایانه نسبت به مساله دارد: مساله به صورت اجزاء رفتاری (و نه کارکردی) تقسیم می‌شود، هر لایه می‌تواند کنترل کل عامل را در اختیار بگیرد و هر کدام از رفتارها نگاشت‌ای مستقیم بین ورودی سنسور و خروجی حرکتی‌ی عامل است. همچنین با وجود آن که لایه‌های بالاتر توانایی محدود کردن لایه‌های پایین را دارند، اما به طور مستقیم فرمان‌ای برای انتخاب یا تغییر حالت لایه‌های پایین نمی‌دهند. به عبارت دیگر معماری‌ی پیش‌نهادی – برخلاف Feudal RL – معماری‌ای گسترده است و به دلیل وجود نداشتن هماهنگ‌کننده‌ی مرکزی نسبت به خطاهای پیش‌آمده مقاوم‌تر است. علاوه بر این‌ها تفاوت حالت حس‌شده‌ی هر رفتار فقط در میزان تجرید آن نیست، بلکه هر رفتار می‌تواند از مجموعه‌ای کاملاً متفاوت از حالت‌های موجود در عامل برای تصمیم‌گیری استفاده کند.

روش MaxQ Value Decomposition – که آن را می‌توان نسخه‌ی پیش‌رفته‌تر و کامل‌تر Feudal RL دانست – یکی دیگر از روش‌های متداول یادگیری تقویتی سلسله‌مراتبی است [Dietterich00]. در این روش وظیفه‌ی عامل به تعدادی زیر-وظیفه‌ی^{۴۷} ساده‌تر – ولی کلی‌تر – تقسیم

⁴⁷ Subtask

می‌شود. هر کدام از این زیر-وظیفه‌ها مانند زیربرنامه‌هایی هستند که توسط قسمت اصلی برنامه‌ای فراخوانی می‌شوند. عامل برای انجام دادن هر وظیفه با توجه به حالت محیط یکی از این زیربرنامه‌ها را فراخوانی می‌کند. مثلاً برای روبات‌ای که می‌بایست جسمی را از قسمتی از اتاق به قسمت‌ای دیگر ببرد، می‌توان چنین زیر-وظیفه‌هایی را تعیین کرد: "حرکت به سمت مکان‌ای مشخص"، "برداشتن جسم"، "قرار دادن جسم" و "دور زدن موانع". عامل در بالاترین سطح تصمیم‌گیری با توجه به حالت فعلی عامل به‌ترین زیر-وظیفه را انتخاب می‌کند. هر کدام از زیر-وظیفه‌ها خود نیز می‌توانند زیر-وظیفه‌های دیگری را فراخوانی کنند. به مانند options، روش MaxQ را می‌توان به صورت SMDP فرمول‌بندی کرد.

روش MaxQ شباهت‌هایی با روش یادگیری ساختار و رفتاری که در این پژوهش معرفی می‌شود دارد. شباهت دو روش در این است که هر دوی آن‌ها ارزش⁴⁸ هر حالت را به صورت ارزش اجزای ساده‌تری می‌نویسند: MaxQ ارزش کلی را بر حسب ارزش زیر-وظیفه‌ها می‌نویسد و روش من ارزش کلی را بر حسب ارزش رفتارها و لایه‌ها می‌نویسد (یادگیری ساختار). این موضوع پس از توضیح روش واضح خواهد شد. از دیدی دیگر می‌توان آن را مشابه با روش یادگیری رفتار دانست: در هر دو روش تخمین‌ای از ارزش زیر-وظیفه‌های پایین‌تر (لایه‌های پایین‌تر در یادگیری رفتار) در هر تصمیم‌گیر نگه‌داری می‌شود. با این‌همه روش MaxQ روش‌ای است که ساختار در آن ثابت است و یادگیری‌ای بر روی ساختار انجام نمی‌شود. در حالی که در روش‌ای که معرفی خواهیم کرد ساختار تغییرپذیر خواهد بود. همچنین روش پیش‌نهادی – همان‌طور که در توضیحات Feudal RL گفته شد – روشی رفتارگرا است در حالی که MaxQ این‌گونه نیست. در نتیجه اگر بخشی از معماری MaxQ برداشته شود، عامل دیگر قادر به ادامه‌ی کار نیست در حالی که با برداشته شدن یا خراب شدن بخشی از معماری رفتارگرا، عامل همچنان به وظیفه‌ی خود عمل می‌کند – گیریم با بازده پایین‌تر.

فرض کنید که کنترلر عامل‌ای به صورت یک برنامه با زبان برنامه‌نویسی متداولی – مانند Lisp – یا یک Finite State Machine (FSM) توصیف شده باشد. بدیهی است که تصمیم‌هایی که عامل اخذ می‌کند بستگی به برنامه‌ای دارد که کنترلر آن عامل را توصیف می‌کند و آن برنامه نیز به

⁴⁸ Value

صورت از پیش توسط برنامه‌نویس مشخص شده است. به عنوان نمونه، طراح تشخیص می‌دهد که برای پوشش دادن⁴⁹ کل یک محوطه می‌بایست همیشه به موازات دیوار سمت راست حرکت کرد و البته گاهی — که با احتمالی مشخص می‌شود — از دیوار کنده شد تا به دیوار دیگری رسید. با وجود سادگی و موثر بودن این روش، عیب آن تطبیق‌ناپذیر بودن کنترلر طراحی شده است. ممکن است سیاست دیگری به عامل‌ای با بازده بالاتر منجر شود و یا این که طراحی‌ای که در محیط قدیم صورت گرفته دیگر چندان مناسب محیط تغییر یافته نباشد. ایده‌ی دو روش Hierarchies of Abstract Machines (HAM) و ALisp بر افزودن قابلیت یادگیری بر کنترلری این چنین توصیف شده است.

در HAM از توصیف FSM⁵⁰ی استفاده می‌شود که در گره‌های انتخاب‌اش⁵¹ امکان یادگیری اضافه شده است ([Parr98a] و [Parr98b]). عامل با توجه به سیگنال تقویت به‌ترین انتخاب از بین چندین انتخاب ممکن در هر چنان گره‌ای را یاد می‌گیرد. این روش که تئوری ریاضی‌اش — درست مانند options و MaxQ — بر پایه‌ی SMDP است با مقید کردن سیاست‌های تحقق‌پذیر، ابعاد فضای فرضیه^{52,51} را کوچک می‌کند. این کاهش ابعاد فضا باعث سریع‌تر شدن یادگیری می‌شود. ALisp گسترش‌ای از HAM است که در آن برخلاف HAM کنترلر با زبان‌ای سطح بالا مانند Lisp توصیف می‌شود ([Andre02] و [Andre03]). گفتنی است روش ALisp از تجزیه تابع ارزش‌ای مشابه با MaxQ استفاده می‌کند.

همان‌طور که اعمال ساختار (برنامه) از پیش مشخص برای FSM (Lisp) باعث کاهش قابل ملاحظه‌ی ابعاد فضای فرضیه می‌شود، به همین صورت مقید کردن نتیجه به سیاست‌های محدود نیز ممکن است باعث پایین آمدن بازده نهایی عامل شود. اگر دانش اولیه‌ی طراح دقیق نباشد و ساختاری (برنامه‌ای) نه چندان خوب طراحی کند وجود یادگیری در نقاط انتخاب نیز باعث رسیدن به پیشینه‌ی بازده ممکن نخواهد شد. به عبارت دیگر بهینه‌گی سیاست نهایی کاملاً به دانش طراح از مساله باز می‌گردد. این خانواده‌ی روش‌ها برخلاف روش‌های یادگیری‌ای که در فصل‌های بعد پیش‌نهاد خواهد شد قابلیت

⁴⁹ Coverage

⁵⁰ Choice Node

⁵¹ Hypothesis Space

⁵² برای دیدن تعریف فضای فرضیه به کتاب‌های یادگیری ماشینی مانند [Mitchell97] یا [Vidyasagar03] مراجعه کنید

طراحی‌ی خودکار سلسله‌مراتب را ندارند و مشابه با یادگیری رفتار -بدون یادگیری ساختار- معرفی شده در این پایان‌نامه‌اند. قابل ذکر است که یادگیری رفتار پژوهش فعلی تا حدی مشابه با یادگیری‌ای است که در این دو روش مطرح می‌شود و به نظر می‌آید که این روش‌ها به دلیل قدرت توصیفی بالایی که دارند قادر به ایجاد معماری‌ای مشابه با معماری‌ای که در این پایان‌نامه به آن می‌پردازم هستند. البته همچنان مشابه با قبل تفاوت در رفتارگرا بودن معماری‌ی این پژوهش با بقیه‌ی معماری‌ها و روش‌های یادگیری تقویتی سلسله‌مراتبی به قوت خود باقی است.

(۲) معرفی مساله

در این فصل بستر ریاضی روش‌های مورد بحث در این پایان‌نامه را معرفی می‌کنیم. با استفاده از این بستر ریاضی می‌توان به توسعه‌ی یادگیری ساختار، یادگیری رفتار و تکامل رفتار پرداخت.

فرض کنید که مجموعه‌ای از n رفتار مختلف در اختیار داریم. این مجموعه را با $\{B_i\} (i=1, \dots, n)$ نمایش می‌دهیم. هر رفتار را -که نگاشت‌ای بین ادراک عامل و عمل آن است- بدین‌گونه توصیف می‌کنیم:

$$\begin{aligned}
 B_i : S'_i &\rightarrow A'_i \quad i=1, \dots, n \\
 A'_i &= A_i \cup \{\text{No Action}\} \\
 S'_i &= \{s'_i | s'_i = M_i(s); \forall s \in S_i\} \\
 S_i &\subset S, A_i \subset A \\
 M_i : S &\rightarrow S'_i
 \end{aligned} \tag{۲-۱}$$

که در آن S بیان‌گر حالت کلی عامل است و A نیز مجموعه‌ی اعمالی است که عامل می‌تواند با استفاده از آن محرک‌های اش^۱ را تغییر دهد. مثلاً در یک روبات متحرک، S معادل اطلاعات مقادیر خوانده شده توسط فاصله‌یاب‌های سونار^۲، سنسورهای برخوردی^۳، اطلاعات دریافتی از دوربین، تخمین از مکان فعلی روبات و ... است. A نیز می‌تواند مجموعه‌ی دستورهای حرکتی‌ای باشد که به موتور چرخ‌ها داده می‌شود. در روابط بالا S_i زیرمجموعه‌ای از کل فضای حالت عامل است که توسط رفتار B_i رویت‌پذیر^۴ است. در حالت کلی این زیرمجموعه برای رفتارهای مختلف یک‌سان نیست ($S_i \neq S_j$) و البته هیچ دلیلی هم وجود ندارد که زیرفضایی که رفتارهای مختلف حس می‌کنند بدون اشتراکی با هم

¹ Actuators

² Sonar

³ Bumper

⁴ Observable

باشند ($S_i \cap S_j \neq \emptyset$). معمولاً هر رفتار تنها به بخش کوچکی از کل فضای حالت عامل دسترسی دارد. مثلاً رفتاری که برای سریع واکنش دادن و جلوگیری از برخورد با اجسام ساخته شده است کافی است به سنسورهای برخوردی و فاصله‌یاب‌ها دسترسی داشته باشد و مثلاً اطلاع از این که در فاصله‌ی دور جسمی قرار دارد معمولاً لازم نیست. مشابه با قبل A_i نیز مجموعه‌ی اعمالی است که رفتار B_i می‌تواند انجام دهد؛ انتخاب عضوی از A_i باعث تغییر وضعیت یکی از خروجی‌های عامل می‌شود. M_i نیز نگاشت‌ای است که حالت کلی‌ی عامل را به حالت درونی‌ای که رفتار B_i برای تصمیم‌گیری‌هایش استفاده می‌کند تبدیل می‌کند. مثلاً M_i می‌تواند معادل حذف تعدادی از حالت‌های S (با فرآیند تصویر کردن^۵ به زیرفضایی از فضای اصلی) یا استخراج تعدادی ویژگی از حالت‌های اصلی باشد.

هر رفتار B_i با عمل‌ای فرضی به نام "عمل پوچ" (یا ناعمل^۶) که آن را با NA نشان خواهیم داد توسعه یافته است. NA عمل‌ای است که حتی در صورت انتخاب نیز باعث تغییر محرک‌ها نمی‌شود.^۷ برای این که ببینید یک رفتار چگونه عمل می‌کند، فرض کنید که عامل در حالت $s \in S$ است. هر کدام از رفتارهای B_i از طریق نگاشت M_i به زیرفضای نگاشت‌یافته‌ی S'_i از S دسترسی دارند. اگر $s_i \in S_i$ (یا $s'_i \in S'_i$) در بازنمایی^۸ داخلی هر رفتار، آن‌گاه رفتار B_i **تحریک**^۹ می‌شود. حال اگر رفتار B_i یکی از اعمال حقیقی‌ی خود و نه NA را انتخاب کند، آن‌گاه به آن **فعال‌شده**^{۱۰} می‌گوییم: اگر $a_i = B_i(M_i(s_i)) = B_i(s'_i) \neq NA$ اما در صورت‌ای که آن رفتار تحریک نشود و یا تحریک بشود ولی عمل NA را انتخاب کند، رفتار هیچ عمل‌ای را بر روی محرک‌های عامل انجام نداده است.

⁵ Projection

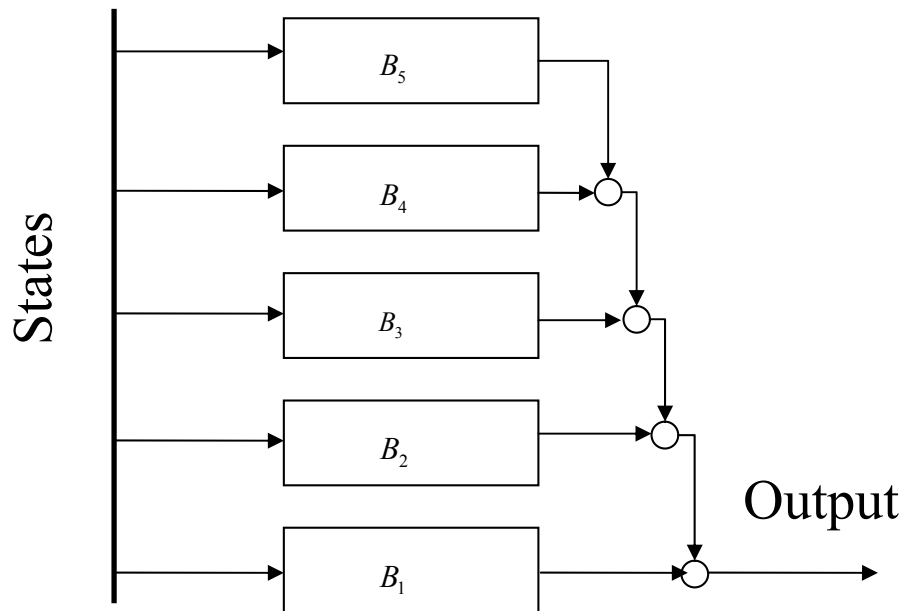
⁶ No Action

⁷ دلیل وجود NA در فصل یادگیری (فصل ۴) رفتار مشخص می‌شود اما به طور خلاصه NA باعث می‌شود تا رفتار گاهی حتی در صورت تحریک شدن نیز خروجی‌ای ندهد و اجازه دهد تا رفتارهای پایین‌تر فعال شوند. بدین‌گونه رفتارهای پایین‌تر ممکن است بتوانند عمل‌ای را پیش‌نهاد دهند که نتیجه‌ی به‌تری برای عامل دارد و بازده بالا رود. این که در چه شرایطی می‌بایست چنین انتخاب‌ای صورت پذیرد در همان فصل بررسی می‌شود.

⁸ Representation

⁹ Excite

¹⁰ Activated



شکل ۱-۲. نمونه‌ای از معماری Subsumption موازی

تا به حال به معرفی ی تک‌رفتار پرداخته بودم. هدف این پژوهش بررسی ی معماری‌های رفتارگرا است. معماری‌هایی که در آن تعدادی رفتار با کنش با هم رفتار کلی ی عامل را مشخص می‌کنند. همان‌طور که پیش‌تر گفته‌ام، هدف این پروژه خودکارسازی ی تولید معماری‌های رفتارگرای سلسله مراتبی است. به همین دلیل، معماری ی انتخاب‌شده حالت خاص ای از معماری ی متداول و مشهور *Subsumption* است. در این حالت خاص که به آن معماری *Subsumption* کاملاً موازی^{۱۱} می‌گوییم، همه ی رفتارها به صورت موازی هم قرار گرفته‌اند (شکل ۱-۲). رفتارهای که در لایه‌های بالایی معماری قرار دارند نسبت به رفتارهای لایه‌های پایین‌تر خود اولویت دارند. اولویت‌شان نیز بدین صورت است که در صورت فعال‌شدن دو رفتار B_1 و B_2 (به خاطر داشته باشید که فعال شدن رفتار معادل پیش‌نهاد خروجی دادن است)، رفتار بالاتر (مثلاً B_2) کنترل عامل را در اختیار قرار می‌گیرد و اجازه ی در اختیار داشتن کنترل عامل را به رفتار B_1 نمی‌دهد و خود به تغییر محرک‌های عامل دست

^{۱۱} Purely Parallel SubSumption Architecture (PPSSA)

می‌یازد. در این حالت گفته می‌شود که B_2 رفتار **کنترل‌کننده**^{۱۲} در معماری است و B_1 نیز توسط B_2 **سرکوب**^{۱۳} شده است.

اینک می‌خواهم نمایش لازم برای معرفی یک معماری PPSSA را معرفی کنم. فرض کنید که n رفتار $\{B_i\}$ داریم و قصد بیان معماری m -رفتاری T را که شامل m تا از کل رفتارها است داریم ($m \leq n$). $T(i)$ نمایان‌گر رفتار لایه‌ی i ام معماری T است. در این نوشتار، شماره‌گذاری از پایین‌ترین لایه آغاز می‌شود: $T(1)$ پایین‌ترین رفتار معماری است و $T(m)$ بالاترین آن. به عبارت دیگر داریم:

$$T = [B_{index(1)} \ B_{index(2)} \ \dots \ B_{index(m)}]^T \quad m \leq n \quad (2-2)$$

$index(i): j$ (that indicates B_j is in the i^{th} layer)

پس مثلاً معماری

$$T = [Wandering \ \text{BallCollection} \ \text{ObstacleAvoidance}] \quad (2-3)$$

بیان‌گر ساختار معماری عامل‌ای است که در پایین‌ترین و کم اولویت‌ترین لایه رفتار پرتاب قرار گرفته است و در بالاترین و پر اولویت‌ترین لایه رفتار جلوگیری از برخورد.

اکنون به مساله‌ی طراحی خودکار عامل هوشمند می‌پردازم. همان‌گونه که پیش‌تر گفتم این مساله معادل با پیدا کردن طرحی‌ای است که معیار مشخص شده توسط طراح را بهینه کند. برای یافتن این بهینه‌گی باید جستجویی در فضای طرح‌های ممکن انجام داد. هدف طراحی نیز پیدا کردن عضوی از فضای بزرگ طرح‌ها است که آن معیار عینی را بهینه کند. پس برای عملی کردن این کار، می‌بایست هر طرح را به صورت پارامترهایی در فضای پارامترها توصیف کرد. از دیدگاهی دو مساله مطرح می‌شود:

¹² Controlling

¹³ Suppress

- هر طرح ممکن به چه صورت بازنمایی می‌شود؟
- جستجو بین طرح‌های مختلف چگونه انجام می‌شود؟

سوال اول به این موضوع می‌پردازد که پاسخ مساله به چه صورت باید توصیف یا بازنمایی شود. به عنوان مثال در بسیاری از موارد از شبکه‌ی عصبی‌ای به عنوان کنترلر عامل هوشمند استفاده می‌کنند و در نتیجه فضای طرح‌ها که توسط پارامترهای شبکه‌ی عصبی -مثلا وزن‌های آن- مشخص می‌شود معادل همه‌ی شبکه‌ی عصبی‌هایی است که با استفاده از معماری‌ی آن شبکه قابل ایجاد است. در این صورت هدف نیز پیدا کردن مجموعه‌ای از وزن‌ها (و یا پارامترهای دیگر) است تا معیار مورد نظر بهینه شود. در این حالت بازنمایی به صورت شبکه‌ی عصبی بوده است.

برای جستجو بین طرح‌های گوناگون روش‌های مختلفی وجود دارد. در این پژوهش از دو روش بهینه‌سازی یادگیری تقویتی و تکامل مصنوعی (یا تکامل به طور خلاصه) استفاده می‌کنم. البته فرض‌های پشت هر کدام از این روش‌ها آنقدر متفاوت است که نگاه کردن به آن دو تنها به عنوان روش‌ای برای جستجو، ساده‌سازی نادقیقی است.

ابتدا به یادگیری تقویتی و کاربردش در این پژوهش می‌پردازم: یادگیری تقویتی یکی از روش‌های بسیار متداول برای تطبیق موجود هوشمند با محیط است. یادگیری یکی از مشخصه‌های عامل هوشمند^{۱۴} است. عامل هوشمند می‌بایست برای فعالیت در محیط^{۱۵} ناشناس و تغییرپذیر با زمان خود را با رخ داده‌های ناشناخته و هم‌چنین تغییرات محیط تطبیق دهد. چنین عامل‌ای می‌باید با توجه به رفتارش در محیط و پاداش یا تنبیه‌ای که از طرف محیط (یا از امیال درونی‌اش) دریافت می‌کند در مورد رفتارهای آینده‌اش تصمیم‌گیری کند. چارچوب یادگیری تقویتی -یادگیری‌ای که عامل توسط سیگنال تقویت رفتار آتی‌اش را مشخص می‌کند- مدت‌هاست که در روان‌شناسی به عنوان توصیف‌گر یادگیری موجودات مطرح شده است [Hergenhahn01]. خوش‌بختانه این دیدگاه دارای فرمول‌بندی ریاضی‌ی دقیقی نیز هست و در سال‌های اخیر تلاش‌های بسیاری برای به کارگیری آن در عامل‌های هوشمند

¹⁴ Intelligent Agent

¹⁵ Environment

ماشینی شده است ([Sutton88], [Watkins89], [Watkins92], [Jaakkola94], [Singh00]). برای اطلاعات بیشتر و توضیح کلی در مورد این روش‌ها به [Kaelbling96] و [Sutton98] و همچنین منابع معرفی شده در آن‌ها مراجعه کنید. برای توصیف روش‌های متداول یادگیری تقویتی سلسله مراتبی - که شباهت‌هایی با نتایج یادگیری این پژوهش دارند - به بخش یادگیری تقویتی سلسله مراتبی (فصل ۱) نگاه کنید.

اگر r_t را به عنوان سیگنال تقویت^{۱۶} دریافتی (پاداش یا تنبیه) در زمان t بدانیم (که می‌تواند به حالت سیستم، عمل انتخابی، حالت جدید و ... بستگی داشته باشد)، آن‌گاه می‌توان پاداش بازگشته^{۱۸،۱۷} را به یکی از صورت‌های زیر نوشت:

$$R = \frac{1}{N} \sum_{i=1}^N r_i \quad (۲-۴)$$

$$R = \sum_{i=1}^{\infty} \gamma^i r_i \quad (0 \leq \gamma < 1) \quad (۲-۵)$$

$$R = \lim_{N \rightarrow \infty} \left(\frac{1}{N} \sum_{i=1}^N r_i \right) \quad (۲-۶)$$

که در این‌جا تنها به مورد اول یعنی میزان پاداش در یک اپیزود می‌پردازم. در این صورت می‌توان ارزش کل معماری با ساختار T و مجموعه رفتارهای $\{B_i\}$ ($i = 1, \dots, n$) را بدین صورت نوشت:

$$\begin{aligned} V_T &= E_{\pi} \left[\frac{1}{N} \sum_{i=1}^N r_i \mid \text{the agent with structure } T \text{ and set of behaviors } \{B_i\} (i = 1, \dots, n) \right] \\ &= E_{\pi} [R \mid \text{the agent with structure } T \text{ and set of behaviors } \{B_i\} (i = 1, \dots, n)] \end{aligned} \quad (۲-۷)$$

^{۱۶} Reinforcement Signal

^{۱۷} Return

^{۱۸} انتخاب پاداش بازگشته به جای بازگشت به دلیل جلوگیری از گنگ شدن کلمه‌ی متداول بازگشت بوده است.

که امید ریاضی بر روی همه‌ی حالت‌های سیستم در طول دوره‌ی زندگی عامل حساب می‌شود و توزیع احتمال آن نیز بر اساس سیاست عامل (π) و توزیع احتمال انتقال حالت^{۱۹} مشخص می‌شود. همان‌طور که بعدتر خواهید دید، سیاست عامل به سیاست هر کدام از رفتارها و ساختار آن رفتارها در معماری عامل بستگی دارد و در نتیجه این سیاست کل، تابع‌ای از سیاست‌های جزئی‌تر است. در جاهای دیگر این پایان‌نامه ممکن است این ارتباط امید به سیاست (π) به طور صریح ذکر نشود.

هدف یادگیری، بیشینه‌گی V_T است. چنین کاری را در دو سطح می‌توان انجام داد: سطح ساختار و سطح رفتار.

در یادگیری ساختار معماری فرض می‌کنیم که مجموعه‌ای مناسب از رفتارها $(\{B_i\} \mid i=1, \dots, n)$ وجود دارد و هدف انتخاب m رفتار از n رفتار موجود و مرتب کردن آن‌ها در معماری‌ای (که به صورت T نمایش داده می‌شود) است به صورتی که (۷-۲) بیشینه شود:

$$T^* = \arg \max_T V_T. \quad (۲-۸)$$

فرض وجود مجموعه‌ی رفتارهای مناسب چندان غیرواقعی نیست. در بسیار از موارد طراح از این که برای حاصل شدن رفتار مورد نظر، چه رفتارهای ساده‌تری لازم‌اند تا حدی آگاه است ولی نمی‌داند هر کدام از آن رفتارها چه اولویتی می‌بایست داشته باشند. پس او می‌تواند مجموعه‌ای از رفتارها را طراحی کند و آن را به صورت جعبه‌ابزاری در اختیار عامل یادگیرنده قرار دهد. این مساله در فصل ۳ بررسی می‌شود.

از سویی دیگر، مساله‌ی یادگیری رفتارها – به عنوان دیگر پارامتر توصیف‌کننده معماری – مطرح می‌شود. هدف یادگیری رفتار پیدا کردن نگاشت‌ای بین فضای حس‌شونده توسط هر کدام از رفتارها (S'_i) به فضای اعمال ممکن برای آن رفتار (A'_i) است. به عبارت دیگر عامل می‌بایست با فرض وجود ساختاری ثابت (ترتیب رفتارها در معماری) یاد بگیرد که به ازای هر s_i (یا s'_i از دید هر رفتار) چه a_i^* را انتخاب کند تا میزان پاداش دریافتی‌اش بیشینه شود:

¹⁹ State Transition Probability

$$\{B_i^*\} = \arg \max_{B_i} V_T. \quad (2-9)$$

مطابق آن چه در یادگیری تقویتی موسوم است می بایست بازنمایی ای برای نگاه داری اطلاعات حاصل از تعامل با محیط وجود داشته باشد. این بازنمایی می بایست نشان دهد که ارزش انجام هر عمل در هر حالت ممکن چیست. برای انجام این کار برای هر رفتار تابع ارزش حالت-عمل را تعریف می کنیم:

$$Q_i(s'_i, a_i), s'_i \in S'_i, a_i \in A'_i \quad \forall B_i \in \{B_i\} \quad (2-10)$$

همچنین نگاشت زیر بین آن چه هر رفتار حس می کند و عمل اش وجود دارد:

$$B_i(s'_i) = \text{Action Selection Mechanism}(Q_i(s'_i, \cdot)) \quad (2-11)$$

که در هر حالت انتخاب عمل بهینه بدین صورت خواهد بود:

$$B_i(s'_i) = \arg \max_{a_i} Q_i(s'_i, a_i). \quad (2-12)$$

ذکر دو نکته مفید است: نخست آن که در بازنمایی بالا، تابع ارزش-حالت برای مجموعه اعمال توسعه یافته (A'_i) تعریف می شود و نه فقط اعمال ممکن. دلیل آن نیز همان طور که پیش تر به طور مختصر گفته شد و جلوتر دقیق تر بررسی خواهد شد- اجازه دادن به رفتار برای فعال نشدن با وجود تحریک شدن است. با این کار رفتارهای زیرین که ممکن است گزینه های بهتری برای کنترل عامل باشند اجازه کنترل کننده بودن می یابند. نکته ی دوم این که انتخاب عمل معمولاً دارای اندکی رفتار

کاوش گرانه^{۲۰} است و در نتیجه (۷) معمولاً گزینه‌ی به کار گرفته شده نیست و به جای‌اش از مکانیزم‌های انتخاب عمل‌ای مانند ϵ -greedy یا انتخاب عمل بولتزمن^{۲۱} استفاده می‌شود.

می‌توان این دو نوع یادگیری را با هم به کار گرفت؛ جستجوی طرح بهینه را در هر دو فضای پارامترهای رفتار و پارامترهای ساختار انجام داد. در این صورت، هدف یادگیری به شکل زیر بیان پذیر است:

$$\{T^*, B_i^*\} = \arg \max_{T, B_i} V_T. \quad (۲-۱۳)$$

یادگیری رفتار و یادگیری هم‌زمان رفتار و ساختار در فصل ۴ مورد مطالعه قرار می‌گیرد.

استفاده از تکامل مصنوعی نیز راه دیگری برای ایجاد کنترلرهای مناسب است. تکامل مصنوعی که الهام گرفته از فرآیند تکامل طبیعی و اصل انتخاب طبیعی است با ایجاد جمعیت‌ای از راه‌حل‌های ممکن برای مساله‌ی مورد نظر و مقایسه‌ی کیفیت هر کدام از راه‌حل‌ها با هم و استفاده از اپراتورهای تکاملی‌ای مانند تزویج^{۲۲} و جهش^{۲۳} به پاسخ مطلوب نزدیک می‌شود. در این‌جا چندان وارد بحث تکامل نمی‌شوم و خواننده را به نوشته‌های کلاسیکی چون [Goldberg89] ارجاع می‌دهم.

در این پژوهش، از تکامل برای ایجاد رفتارهای مطلوب استفاده شده است (فصل ۵). به عبارت دیگر برای یافتن رفتارهای بهینه اضافه بر استفاده از یادگیری، از تکامل نیز استفاده شده است. بازنمایی رفتارها با تغییرات اندکی مشابه با آن چیزی است که در چند پاراگراف پیش‌تر در بخش یادگیری رفتار توضیح دادم. اما در این‌جا دیگر تابع ارزش-عمل تعریف نمی‌شود و به طور مستقیم سیاست هر رفتار بازنمایی می‌شود. به عبارت دقیق‌تر هر رفتار B_i را به شکل نگاشت‌ای بین ورودی و خروجی آن رفتار مشخص می‌کنم:

²⁰ Exploration

²¹ Boltzman Action Selection

²² Crossover

²³ Mutation

$$B_i(s') : S'_i \rightarrow A'_i. \quad (۲-۱۴)$$

خروجی این نگاشت از میان مجموعه‌ی توسعه‌یافته‌ی اعمال (A'_i) انتخاب می‌شود. در صورتی که عمل انتخابی NA بوده باشد، رفتار فعال نخواهد شد و به رفتارهای پایینی - در صورت وجود - اجازه‌ی فعال شدن می‌دهد. در پیاده‌سازی انجام شده این نگاشت به صورت جدول ارجاعی^{۲۴} با ابعاد مناسب بازنمایی می‌شود در حالی که هیچ محدودیت ویژه‌ای برای نوع بازنمایی وجود ندارد.

در نهایت لازم به تاکید است که جستجوهای انجام شده سعی در پیدا کردن بهینه‌ی مقید می‌کنند. مثلاً به هنگام حل مساله‌ی یادگیری ساختار، جستجو در فضای مقید شده به رفتارهای موجود انجام می‌شود - در نتیجه بدیهی است که مساله ممکن است پاسخ‌های بهتری نیز داشته باشد.

²⁴ Lookup Table

۳) یادگیری ساختار

۳-۱) مقدمه

در دو فصل پیش درباره‌ی روی‌کرد کلی‌ی این پژوهش به فرآیند خودکارسازی طراحی‌ی معماری‌ی رفتارگرا توضیح دادم. در این فصل دو روش متفاوت یادگیری ساختار برای خودکارسازی‌ی طراحی‌ی معرفی می‌کنم.^۱

هدف یادگیری ساختار – همان‌طور که پیش‌تر گفته شد – پیدا کردن ساختار T^* ای است که (۷-۲) را بهینه کند. در یادگیری ساختار فرض بر آن است که مجموعه‌ای از رفتارها از-پیش در اختیار است و هدف پیدا کردن به‌ترین انتخاب و ترتیب از آن‌ها است. یک روش مناسب یادگیری ساختار می‌بایست به سوال‌های زیر پاسخ دهد:

- **بازنمایی:** عامل چگونه می‌بایست دانش کسب شده در دوره‌ی حیات‌اش را ذخیره کند؟
- **تخصیص اعتبار در ساختار سلسله مراتبی:** عامل دارای ساختار سلسله مراتبی چگونه می‌بایست رفتار مسوول در عمل‌کردش را تشخیص داده و به تناسب آن را تشویق یا تنبیه کند؟
- **به روز رسانی دانش:** عامل در مواجهه با تجربه‌ای جدید چگونه می‌بایست دانش قبلی‌اش را به روز رساند؟

قدم اول انتخاب بازنمایی‌ی مناسب‌ای از دانش‌ای است که عامل در طول زندگی‌اش کسب کرده. این بازنمایی می‌بایست آینده‌ی تمام‌نمایی از تجربیات عامل باشد. عامل به کمک این بازنمایی است که قادر به تصمیم‌گیری درباره‌ی به‌ترین ساختار ممکن می‌شود. به عبارت دیگر این بازنمایی باید به عامل اجازه دهد تا مقادیر مختلف V_T را محاسبه کرده و ساختاری را که بالاترین مقدار آن را پیش‌بینی می‌کند انتخاب کند. این بازنمایی لازم است چند ویژگی‌ی مختلف داشته باشد: در ابتدا این

^۱ نتایج مقدماتی‌ی این بخش در [Farahmand04] منتشر شده است.

بازنمایی می‌بایست کامل باشد و بتواند همه‌ی حالت‌های ممکن ساختار را توصیف کند. مثلاً فرض کنید تعداد n رفتار در دسترس است و ما می‌خواهیم معماری‌ای با m رفتار ایجاد کنیم. یک نوع بازنمایی می‌تواند این باشد که عامل امتیازی را که هر رفتار در طول دوره‌ی زندگی‌اش کسب کرده ثبت کند و در مرحله‌ی انتخاب معماری‌ی جدید با توجه به تجربیات پیشین‌اش m رفتار با امتیاز بالاتر را انتخاب کند. این روند باعث ایجاد معماری‌ای می‌شود که به احتمال بالا باعث دریافت متوسط پاداش بیش‌تری در طول عمر عامل نسبت به حالت انتخاب تصادفی معماری‌ها می‌شود. اما این معماری از بین تعدادی محدود از معماری‌های ممکن انتخاب شده است و فضای جستجوی آن بسیار کوچک‌تر از فضای n^m (انتخاب n رفتار برای هر کدام از m لایه) و یا $n!/(n-m)!$ (در نظرگیرنده‌ی ترتیب و بدون تکرار) معماری ممکن است. در نتیجه بسیاری از معماری‌ها در این فضا قابل توصیف نیستند و اولویت رفتارها نسبت به هم به هیچ صورت‌ای مشخص نمی‌شود.

به عنوان مثال‌ای از بازنمایی‌ی کامل می‌توان به بازنمایی‌ای اشاره کرد که در آن امید ریاضی پاداش (یا امتیاز) n^m معماری‌ی ممکن ذخیره شده است. این بازنمایی کامل است ولی دارای فضای بسیار بزرگی است. عامل برای آن که تصمیم بگیرد کدام رفتارها را و با چه ترتیب‌ای انتخاب کند، می‌بایست تجربیات بسیاری زیادی را انجام دهد در حالی که بسیاری از این تجربه‌ها لازم نیستند. این بازنمایی خاص -ولی کامل- استفاده‌ای از ارتباط احتمالی بین تجربه‌های دو معماری‌ی مختلف نمی‌کند و آن‌ها را از نظر منطقی کاملاً مستقل از هم می‌داند در حالی که به عنوان مثال عامل با تجربه کردن این که یک جفت رفتار وقتی ترتیب خاص‌ای داشته باشند همیشه باعث پاداش گرفتن می‌شوند می‌تواند فضای جستجوی‌اش را بسیار کوچک کند. در واقع بازنمایی‌ای مطلوب است که عامل را قادر به تعمیم^۲ دانش‌اش بکند؛ گو این که این تعمیم ممکن است همیشه درست نباشد. پس به طور خلاصه ویژگی‌های یک بازنمایی خوب این‌ها می‌اند:

○ کامل^۳ باشد.

^۲ Generalization

^۳ Complete

○ فضای بازنمایی^۴ کوچک‌ای داشته باشد.^۵

○ از تجربیات، هوش‌مندانه بهره‌برد.

مساله‌ی مهم دیگر برای یادگیری ساختار، تخصیص اعتبار در ساختار سلسله‌مراتبی معماری است. چگونگی اعمال نتیجه‌ی این تخصیص اعتبار بر دانش ذخیره شده در بازنمایی، مساله‌ی به روزرسانی بازنمایی را ایجاد خواهد کرد. انتخاب مناسب نحوه‌ی بازنمایی علاوه بر ویژگی‌های ذکر شده پیشین حل مساله‌ی تخصیص اعتبار سلسله‌مراتبی را ساده می‌کند.

برای حل این سه مساله از ساختار ویژه‌ی معماری Subsumption موازی نهایت استفاده را می‌برم: یعنی در انتخاب نوع بازنمایی به ساختار معماری توجه می‌کنم. برای این کار تلاش می‌کنم تا ساختار را به اجزای ساده‌تری تجزیه^۶ کنم. این تجزیه می‌بایست به گونه‌ای باشد که بازنمایی حاصل شده سه ویژگی ذکر شده در بالا را داشته باشد. سپس با روی‌کردی مرحله به مرحله ابتدا بازنمایی مناسبی را برای ذخیره‌سازی دانش عامل پیش‌نهاد می‌کنم و سپس با توجه به آن بازنمایی نحوه‌ی تخصیص اعتبار و به‌روز رسانی دانش را مشخص می‌کنم. در بخش‌های بعدی دو روش مختلف برای یادگیری ساختار توضیح داده خواهد شد.

۳-۲) یادگیری ساختار با بازنمایی مرتبه صفر

۳-۲-۱) بازنمایی

برای ذخیره‌سازی دانش عامل نسبت به محیط اطراف و تشخیص این‌که چه ترتیبی از رفتارها مناسب‌تر است، از بازنمایی‌ای استفاده می‌کنیم که دانش ما را نسبت به اجزای ساده‌تری از آن معماری

⁴ Representation Space

⁵ در ادبیات یادگیری ماشینی استفاده از فضای فرضیه (hypothesis space) برای چنین تعریفی متداول است. با این حال برای تأکید بر کارکرد بازنمایی این فضا، بیش‌تر از از فضای بازنمایی (representation space) استفاده می‌کنم. برای اطلاع بیش‌تر راجع به تعریف فضای فرضیه می‌توانید به [Mitchell97] و [Vidyasagar03] مراجعه کنید.

⁶ Decompose

ثبت کند. بازنمایی مرتبه صفر - که در این بخش معرفی می‌شود - چنین وظیفه‌ای را بر عهده دارد. در این بازنمایی، **امید ارزش دریافتی هر رفتار در هر لایه** ذخیره می‌شود.

برای روشن شدن ایده‌ی این بازنمایی، فرض کنید که مجموعه‌ای دو تایی از رفتارها $\{B_1, B_2\}$ داشته باشیم و قرار باشد معماری‌ای دو لایه از آن‌ها بسازیم، آن‌گاه - همان‌طور که نشان خواهیم داد - می‌توان با ثبت ارزش‌های زیر، معماری‌ی بهینه را پیدا کرد. فرض کنید که ارزش قرارگیری هر رفتار در هر لایه - که به زودی تعریف دقیق‌ای از آن ارایه خواهیم داد - بدین صورت باشد:

	B_1	B_2
L_1	1	-1
L_2	0.25	0.5

آن‌گاه ارزش معماری‌های ممکن غیرتکراری بدین‌گونه خواهد بود:

$$V_{[B_1 \ B_2]} = 1 + 0.5 = 1.5 \quad (3-1)$$

$$V_{[B_2 \ B_1]} = 0.25 - 1 = -0.75 \quad (3-2)$$

پس ترتیب مناسب $T = [B_1 \ B_2]$ خواهد بود. برای حالت‌های کلی نیز با روندی مشابه می‌توان معماری‌ای یافت که باعث ایجاد با ارزش‌ترین معماری شود (مساله مطرح در رابطه‌ی (۸-۲)). به بیان ریاضی، برای معماری‌ی T می‌توان نوشت:

$$\begin{aligned}
V_T &= E_\pi[R] = E_\pi\left[\frac{1}{N} \sum_{t=1}^N r_t\right] \\
&= E_\pi\left[\frac{1}{N} \sum_{t=1}^N \{r_t \wedge ("L_1 \text{ is controlling"} \vee "L_2 \text{ is controlling"} \vee \dots \vee "L_m \text{ is controlling"})\}\right] \\
&= E_\pi\left[\frac{1}{N} \sum_{t=1}^N \{r_t \wedge "L_1 \text{ is controlling"}\}\right] + \dots + E_\pi\left[\frac{1}{N} \sum_{t=1}^N \{r_t \wedge "L_m \text{ is controlling"}\}\right] \\
&= E_\pi\left[\frac{1}{N} \sum r_t \mid L_1 \text{ is controlling}\right] \cdot P(L_1 \text{ is controlling}) \\
&\quad + E_\pi\left[\frac{1}{N} \sum r_t \mid L_2 \text{ is controlling}\right] \cdot P(L_2 \text{ is controlling}) \\
&\quad + \dots + E_\pi\left[\frac{1}{N} \sum r_t \mid L_m \text{ is controlling}\right] \cdot P(L_m \text{ is controlling})
\end{aligned}
\tag{۳-۳}$$

که در آن $E_\pi\left[\frac{1}{N} \sum r_t \mid L_i \text{ is controlling}\right]$ امید ریاضی دریافت سیگنال تقویت در آن هنگامی است که لایه i ام کنترل عامل را در اختیار گرفته است. $P(L_i \text{ is controlling})$ نیز احتمال در کنترل لایه i ام بودن عامل است. توجه کنید که این تجزیه به دلیل دو به دو مجزا بودن^۷ بودن وقایع $\{t \mid "L_i \text{ is controlling"} \text{ in timestep } t\}$ ممکن شده: در هر زمان فقط یکی از لایه‌ها کنترل عامل را در اختیار دارد. در این فرمول‌بندی فرض شده است که حداقل یکی از لایه‌ها کنترل عامل را در اختیار گرفته‌اند.

با تعریف $V_{zo}(i, j)$ -ارزش مرتبه صفر- به صورت زیر

$$V_{zo}(i, j) = V_{ij} = E_\pi\left[\frac{1}{N} \sum r_t \mid B_j \text{ is the controlling behavior in the } i^{th} \text{ layer}\right] \tag{۳-۴}$$

خواهیم داشت:

⁷ Mutually Exclusive

$$\begin{aligned}
E_{\pi} \left[\frac{1}{N} \sum r_i \mid L_i \text{ is controlling} \right] &= \sum_{j=1}^n P\{B_j \mid L_i\} E_{\pi} \left[\frac{1}{N} \sum r_i \mid B_j \text{ is the controlling behavior in } L_i \right] \\
&= \sum_{j=1}^n P\{B_j \mid L_i\} V_{ij} \quad i=1, \dots, m
\end{aligned}
\tag{۳-۵}$$

که در آن $P\{B_j \mid L_i\}$ احتمال در اختیار B_j بودن کنترل عامل است مشروط به این که کنترل عامل در اختیار لایه‌ی L_i باشد. با در نظر گرفتن روابط به دست آمده، تابع ارزش کل معماری را می‌توان به صورت زیر نوشت:

$$V_T = \sum_{i=1}^m \sum_{j=1}^n P\{B_j \mid L_i\} V_{ij} P(L_i \text{ is controlling}). \tag{۳-۶}$$

مشخص است که این بازنمایی کامل است چون قادر به بازنمایی همه‌ی ترکیب‌های ممکن از رفتارها است. از طرف دیگر ابعاد فضای بازنمایی در حالت بازنمایی مرتبه صفر بسیار کوچک‌تر از فضای بازنمایی کامل‌ای مانند آن که همه‌ی ترکیب‌های ممکن را ذخیره می‌کند است. ابعاد آن فضا برابر با مقدار زیر است:

$$\text{cardinality}(ZO) = n \cdot m. \tag{۳-۷}$$

با داشتن ارزش هر رفتار در هر لایه می‌توان ارزش معماری‌های مختلف را محاسبه کرد و دیگر نیازی به تخمین جداگانه‌ی ارزش هر معماری نیست. این ویژگی باعث می‌شود بدون نیاز به تخمین همه‌ی معماری‌های مختلف، تخمین‌ای از ارزش هر معماری با داشتن تخمین‌ای از ارزش هر رفتار در هر لایه به دست آوریم. بدین‌گونه می‌توان با مقایسه‌ی چندین ساختار فرضی مختلف، به‌ترین آن‌ها را یافت. با توجه به روابط (۲-۸) و (۳-۶)، برای یافتن ساختار بهینه کافی است مساله‌ی زیر حل شود:

$$T^* = \arg \max_T V_T = \arg \max_T \sum_{i=1}^m \sum_{j=1}^n P\{B_j | L_i\} V_{ij} P(L_i \text{ is controlling}). \quad (3-8)$$

برای حل این مساله می‌بایست تخمین‌ای از هر کدام از اجرای آن یعنی V_{ij} ، $P(L_i \text{ is controlling})$ و $P\{B_j | L_i\}$ وجود داشته باشد. در زیربخش بعد نشان خواهیم داد که همه‌ی این مقادیر در طول یادگیری قابل تخمین‌اند.

در یک معماری رفتارگرای متداول در هر لایه فقط یک رفتار وجود دارد. اما در معماری پیش‌نهادی می‌توان با احتمال‌های متفاوتی رفتارهای مختلف را در یک لایه‌ی ثابت انتخاب کرد. به عبارت دیگر یک معماری دیگر ساختاری ثابت نیست بلکه به وسیله‌ی $P\{B_j | L_i\}$ ‌هایی توصیف می‌شود که اجازه‌ی وجود چندین رفتار در یک لایه را به عامل می‌دهد.

حال اگر از چنین خاصیت‌ای استفاده نکنیم و همیشه تنها یکی از رفتارها را در هر لایه انتخاب کنیم - و به عبارت دیگر انتخاب معماری‌مان حریصانه^۸ باشد - آن‌گاه داریم:

$$P\{B_j | L_i\} = \begin{cases} 1 & B_j \text{ is the controlling behavior in } L_i \\ 0 & \text{otherwise} \end{cases} \quad (3-9)$$

و در نتیجه می‌توان نوشت:

$$T^* = \arg \max_T \sum_{i=1}^m V_{i, \text{index}(T(i))} P(L_i \text{ is controlling}). \quad (3-10)$$

۳-۲-۲) تخصیص اعتبار و به روز رسانی بازنمایی

در زیربخش پیشین با بازنمایی مرتبه صفر دانش آشنا شدید. هنگامی که عامل پاداش یا تنبیه‌ای دریافت می‌کند اطلاعات ثبت شده در آن بازنمایی می‌بایست تغییر کند. مطابق با تعریف V_{ij} (۳-۴)، روش

^۸ Greedy

تخصیص اعتبار مشخص است: V_{ij} باید آن هنگام که B_j کنترل لایه i ام را در اختیار داشته باشد به روز شود. پس اگر لایه i ام، لایه i کنترل کننده B_j عامل باشد و B_j در آن لایه فعال شده باشد و در آن هنگام به دلیل عمل انتخاب شده توسط رفتار B_j در آن لایه- عامل سیگنال تقویت r_k را دریافت کند، ارزش V_{ij} مطابق با آنچه در تخمین پارامترهای تصادفی معمول است به روز می شود ([Robbins51] و [Watkins92]):

$$V_{ij,k+1} = (1 - \alpha_{k,ij}) V_{ij,k} + \alpha_{k,ij} r_k \quad (\text{whenever } B_j \text{ is the controlling behavior in the } i^{\text{th}} \text{ layer}) \quad (3-11)$$

البته به دلیل آن که اثباتی از جنس نشان دادن جمع شونده بودن⁹ خطای بین V_{ij} و V_{ij}^* (بهینه V_{ij} - البته در صورت وجود) انجام نداده ام، نمی توان شروط هم گرایی ای مانند آنچه در [Jaakkola94] و موارد مشابه انجام شده است ارایه داد.

تخمین $P(L_i \text{ is controlling})$ و $P\{B_j | L_i\}$ پیچیده نیست. برای تخمین $P(L_i \text{ is controlling})$ می توان شمارنده ای برای هر لایه در نظر گرفت و هر باری که آن لایه کنترل کننده می شود آن را افزایش داد؛ یعنی:

$$P(L_i \text{ is controlling})_{n+1} = \begin{cases} \frac{n P(L_i \text{ is controlling})_n + 1}{n + 1} & \text{if } L_i \text{ is controlling at time step } n \\ \frac{n P(L_i \text{ is controlling})_n}{n + 1} & \text{if } L_i \text{ is not controlling at time step } n \end{cases} \quad (3-12)$$

یا این که از روش تخمین ای مبتنی بر روش روبینز-مونرو ([Robbins51]) که در Q-Learning متداول است ([Watkins89] و [Watkins92]) - استفاده کرد¹⁰:

⁹ Contraction

$$P(L_i \text{ is controlling})_{n+1} = (1 - \alpha_{n,i})P(L_i \text{ is controlling})_n + \alpha_{n,i}["L_i \text{ is controlling in time step } n"] \quad (3-13)$$

توجه کنید که این به روز رسانی در هر واحد زمانی عمومی عامل n انجام می شود و نه هر بار که رفتار فعال می شود. به روز رسانی $P\{B_j | L_i\}$ نیز مشابه است: یا از شمارنده ای استفاده می کنیم یا از روش تخمین روبینز-مونرو-مانند استفاده می کنیم:

$$P\{B_j | L_i\}_{k+1} = (1 - \alpha_{k,i})P\{B_j | L_i\}_k + \alpha_{k,i}["B_j \text{ is active in step } k"] \quad (\text{whenever } L_i \text{ is controlling}) \quad (3-14)$$

و این به روز رسانی هر زمان که لایه L_i کنترل کننده است انجام می شود.¹⁰ در صورتی که $\{\alpha\}$ متناسب با $1/n$ (یا $1/k$ در (3-14)) کاهش یابد، این تخمین ها معادل تخمین متوسط - و با در نظرگیری کل پیشینه ای عامل با وزن یکسان - خواهند بود. با تغییر نحوه ی کاهش $\{\alpha\}$ می توان حافظه ای عامل نسبت به وقایع را نزدیک بین یا دور بین کرد. توجه کنید که مقادیر $\{\alpha_{k,i}\}$ ، $\{\alpha_{n,i}\}$ و $\{\alpha_{k,i}\}$ الزاماً یکسان نیستند.

روش دیگری برای تخمین همزمان همه ی این مقادیر - یعنی V_{ij} ، $P(L_i \text{ is controlling})$ و $P\{B_j | L_i\}$ - وجود دارد. به جای تخمین هر کدام از اجزای (3-6) به طور مجزا، با تعریف $\tilde{V}_{ij}$ به صورت زیر می توان آن ها را همزمان تخمین زد:

$$\tilde{V}_{ij} = P\{B_j | L_i\}V_{ij}P(L_i \text{ is controlling}). \quad (3-15)$$

¹⁰ "Expression" = $\begin{cases} 1 & \text{Expression is True} \\ 0 & \text{Expression is False} \end{cases}$

¹¹ به اختلاف بین این رابطه (3-14) و روابط به روز رسانی $P(L_i \text{ is controlling})$ (3-12) و (3-13) توجه کنید.

در آن صورت خواهیم داشت^{۱۲}:

$$\tilde{V}_{ij,n+1} = (1 - \alpha_{n,ij})\tilde{V}_{n,ij} + \alpha_{n,ij} [B_j \text{ is active at time step } n \times L_i \text{ is controlling at time step } n \times r_n] \quad (3-16)$$

در نتیجه با استفاده از یکی از این دو روش تخمین زدن مقادیر لازم برای (۳-۶) می‌توان آرایش بهینه‌ای از رفتارها را (که (۲-۸) را حل می‌کند) یافت. روش‌های مختلفی برای حل این مساله‌ی بهینه‌سازی ترکیبیاتی^{۱۳} وجود دارد (استفاده از جستجوی کامل فضا^{۱۴}، استفاده از انواع روش‌های سردستی و heuristicها، استفاده از meta-heuristicهایی مانند انجماد شبیه‌سازی شده^{۱۵}، الگوریتم‌های تکاملی، یا شبکه‌های عصبی هاپفیلد^{۱۶} و ...) اما در پیاده‌سازی‌های این پژوهش از جستجوی تصادفی^{۱۷}ی تک نقطه‌ای فضا برای پیدا کردن بهینه آن مساله استفاده شده است.

۳-۳) یادگیری ساختار با بازنمایی مرتبه یک

۳-۳-۱) بازنمایی

در بازنمایی مرتبه صفر - که در بخش پیش معرفی شد - کمیت نگه‌دارنده‌ی دانش، ارزش قرارگیری هر رفتار در هر لایه بود. برخلاف آن در بازنمایی مرتبه یک، ارزش ترتیب رفتارها در معماری مورد توجه قرار می‌گیرد. در این بازنمایی، امید ارزش دریافتی ترتیب‌های دو تایی رفتارها در معماری ذخیره می‌شود. این بازنمایی کامل است و فضای بازنمایی کوچک‌تری نسبت به بازنمایی‌ای که همه‌ی حالت‌های ممکن یک معماری را در نظر می‌گیرد دارد. علاوه بر آن این بازنمایی از ساختار ویژه‌ی

¹² توجه به تفاوت بین این رابطه‌ی جدید (۳-۱۶) و رابطه‌ی (۳-۱۱) که در آن V_{ij} تنها هنگامی به روز می‌شد که B_j رفتار کنترل‌کننده در لایه‌ی i نام بود ضروری است.

¹³ Combinatorial Optimization Problem

¹⁴ Exhaustive Search

¹⁵ Simulated Annealing

¹⁶ Hopfield

¹⁷ Stochastic Search

معماری سلسله مراتبی مان که در آن رفتارها نسبت به هم ترتیب دارند - بالاتر یا پایین تر - بهره می گیرد و روند تقسیم امتیاز را راحت می کند. برای دقیق تر شدن بحث، ابتدا اپراتور ترتیب ' $\succ^T$ ' را بدین گونه تعریف می کنیم:

$$B_i \succ_{B_j}^T B_j : B_i \text{ is upper than } B_j \text{ in the architecture } T|_{B_i, B_j \in T} \quad (3-17)$$

این اپراتور در معماری ای با ساختار T بیان گر بالا یا پایین بودن یک رفتار نسبت به رفتار دیگر است. برای روشن شدن هدف پشت این بازنمایی، مثال ابتدای بخش بازنمایی مرتبه صفر را به گونه ای دیگر می نگریم. فرض کنید که مجموعه ای دو تایی از رفتارها $\{B_1, B_2\}$ داشته باشیم و قرار باشد معماری ای دو لایه از آن ها بسازیم، آن گاه - همان طور که نشان خواهیم داد - می توان با ثبت ارزش های زیر، معماری بهینه را پیدا کرد. فرض کنید که ارزش ترتیب های مختلف B_1 و B_2 این گونه باشند:

Order	Value
$B_1 > B_1$	1.25
$B_1 > B_2$	-0.75
$B_2 > B_1$	1.5
$B_2 > B_2$	-0.5

آن گاه به ترین انتخاب برای ترتیب دو رفتار، انتخاب $B_2 > B_1$ و یا معماری $T = [B_1 \ B_2]$ خواهد بود.

همانند بخش پیش می توان ارزش کل عامل V_T را به صورت ارزش اجزای ساده تری تجزیه کرد:

$$V_T = E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \right] = \sum_{i=1}^m E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_{index(i)[=k]} \text{ is controlling} \right] P(B_{index(i)} \text{ is controlling})$$

(۳-۱۸)

که هر کدام از ارزش‌های عبارت سمت راست بدین صورت قابل بسط است^{۱۸}:

$$\begin{aligned} & E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_k \text{ is controlling} \right] \\ &= E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_k \text{ is controlling and nobody else is active} \right] \\ &+ \sum_{\forall B_j: B_k \succ B_j} E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_k \text{ is controlling and } B_j \text{ is the next active behavior} \right] = V_{k0} + \sum_{\forall B_j: B_k \succ B_j} V_{k \succ j} \end{aligned}$$

(۳-۱۹)

که در آن

$$V_{FO}(B_i \succ B_j) = V_{i \succ j} = E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_i \text{ is controlling and } B_j \text{ is the next active behavior} \right]$$

(۳-۲۰)

$$V_{i0} = E_{\pi} \left[\frac{1}{N} \sum_{t=1}^N r_t \mid B_i \text{ is controlling and nobody else is active} \right].$$

(۳-۲۱)

$V_{i0} (i=1, \dots, n)$ و $V_{i \succ j} (i, j=1, \dots, n)$ اجزای پایه‌ای بازنمایی مرتبه یک هستند. لازمه‌ی معتبر بودن تجزیه‌ی انجام گرفته، دو به دو مجزا بودن جمله‌های (۳-۱۹) است که چنان چیزی از نحوه‌ی

¹⁸ B_j رفتار فعال بعد از B_i (the next active behavior) به صورت اولین رفتار فعال پایین‌تر از B_i ($B_i \succ B_j$) در معماری T تعریف می‌شود.

تعریف V_{i0} و $V_{i>j}$ مشخص است^{۱۹}. با فرض آن که حداقل یکی از رفتارهای معماری فعال شده باشد، آن گاه ارزش عامل را می توان بدین گونه نوشت:

$$V_T = \sum_{i=1}^m \left\{ V_{index(i)0} + \sum_{j=1}^{i-1} V_{index(i)>index(j)} \right\} P(B_i \text{ is controlling}). \quad (3-22)$$

به ترین معماری، آن ای است که (۳-۲۲) را بیشینه می کند. با داشتن توابع ارزش و مقادیر $P(B_i \text{ is controlling})$ ، این مساله ی بهینه سازی ترکیبیاتی با استفاده از یکی از روش هایی که در انتهای زیربخش بازنمایی مرتبه صفر گفته شد قابل حل است. ابعاد بازنمایی مرتبه این چنین است که نسبت به ابعاد فضای کامل بسیار کوچک تر است:

$$cardinality(FO) = n^2 + n. \quad (3-23)$$

۳-۳-۲) تخصیص اعتبار و به روز رسانی بازنمایی

با مشاهده ی الگوی فعالیت رفتارها و با توجه به تعریف V_{i0} و $V_{i>j}$ می توان مساله ی تخصیص اعتبار را حل کرد. در این بازنمایی، دو دسته رخداد فعال شدن یک رفتار و فعال شدن بیش از یک رفتار-پیش می آید. در حالت ای که فقط یک رفتار فعال می شود، V_{i0} می بایست به روز شود. در صورت ای که تعداد دو یا بیش تر رفتار فعال شدند، $V_{i>j}$ به ازای i ای که شماره ی رفتار کنترل کننده B_i و j ای که شماره ی رفتار فعال بعدی B_j است می بایست به روز شود. قوانین به روز رسانی بدین گونه اند:

¹⁹ می توان نشان داد که اگر در تعریف $V_{i>j}$ بقیه ی رفتارهای فعال نیز در نظر گرفته می شد (برخلاف آن چه در این جا به کار رفته است، یعنی فقط اولین رفتار فعال بعدی)، آن گاه اجزای تجزیه ارزش عامل دیگر دو به دو مجزا نبودند و V_T دیگر تجزیه پذیر نبود.

$$V_{i0_{k+1}} = (1 - \alpha_{k,i0})V_{i0_k} + \alpha_{k,i0}r_k \quad \left(\begin{array}{l} \text{whenever } B_i \text{ is controlling} \\ \text{and no other behavior is active} \end{array} \right) \quad (3-24)$$

$$V_{i>j_{k+1}} = (1 - \alpha_{k,i>j})V_{i>j_k} + \alpha_{k,i>j}r_k^T \quad \left(\begin{array}{l} \text{whenever } B_i \text{ is controlling} \wedge B_i^T \succ B_j \\ \wedge B_j \text{ is the next active behavior} \end{array} \right) \quad (3-25)$$

برای تخمین $P(B_i \text{ is controlling})$ می‌توان تعداد مرتبه‌ای را که B_i کنترل کننده بوده است

شمرد. اما به مانند قبل می‌توان هر دو جمله‌ی ارزش و احتمال را با هم به روز رساند. با تعریف

$$\tilde{V}_{i>j} = V_{i>j}P(B_i \text{ is controlling}) \quad (3-26)$$

$$\tilde{V}_{i0} = V_{i0}P(B_i \text{ is controlling}) \quad (3-27)$$

و با استدلال‌ای مشابه بازنمایی پیشین، می‌توان قوانین به روز رسانی را بدین گونه نوشت:

$$\tilde{V}_{i0_{n+1}} = (1 - \alpha_{n,i0})\tilde{V}_{i0_n} + \alpha_{n,i0}r'_{n,i} \quad (3-28)$$

$$\tilde{V}_{i>j_{n+1}} = (1 - \alpha_{n,i>j})\tilde{V}_{i>j_n} + \alpha_{n,i>j}r''_{n,i>j} \quad (3-29)$$

که در آن‌ها

$$r'_{n,i} = \begin{cases} r_n & B_i \text{ is the controlling behavior and no other behavior is active} \\ 0 & \text{otherwise} \end{cases} \quad (3-30)$$

$$r''_{n,i>j} = \begin{cases} r_n & B_i \text{ is the controlling behavior and } B_j \text{ is the next active behavior} \\ 0 & \text{otherwise} \end{cases} \quad (3-31)$$

توجه داشته باشید که (۳-۲۸) و (۳-۲۹) در هر واحد زمانی به روز خواهند شد در حالی که (۳-۲۴) و (۳-۲۵) تنها هنگامی به روز می‌شوند که B_j رفتار کنترل‌کننده باشد.

۴) یادگیری رفتار

۴-۱) مقدمه

در فصل یادگیری ساختار (فصل ۳) فرض کرده بودم که جعبه‌ابزاری از رفتارهای مناسب در اختیار داریم و عامل می‌بایست با استفاده از رفتارهای از پیش تعیین‌شده، به‌ترین ساختار ممکن را بیابد. گاهی طراح دانش نسبی‌ای به مساله‌ی دنیای واقعی دارد و می‌تواند رفتارهایی را که باعث حل مساله می‌شوند تا حدی مشخص و طراحی کند. در این صورت، او می‌تواند از روش‌های معرفی‌شده‌ی یادگیری ساختار برای پیدا کردن به‌ترین ترکیب ممکن رفتارها بهره‌جوید. اما در بیش‌تر موارد دانش طراح به اندازه‌ی کافی و دقیق نیست تا بتواند جزییات هر کدام از رفتارهای مطلوب را مشخص کند و یا حتی اگر بتواند چنان کاری را انجام دهد، هیچ الزامی بر بهینه‌گی طراحی وجود نخواهد داشت. او با توجه به شناختی که از مساله دارد سعی در پیدا کردن به‌ترین نگاشت سنسور-محرك دارد در حالی که این نگاشت بهینه کاملاً به دینامیک کنش‌گرانه‌ی عامل-محیط بستگی دارد و تجربه‌ی بشری طراح معمولاً به یافتن مناسب‌ترین نگاشت منجر نمی‌شود.

برای روشن شدن مسله فرض کنید که طراح قصد دارد تا رفتار حرکتی خاصی را -مثلاً جلوگیری از برخورد با اجسام- به صورت رفتاری برای روبات متحرک‌ای که از سونار استفاده می‌کند ایجاد کند. تجربه‌ی مستقیم او از عمل جلوگیری از برخورد از سنسور قدرت‌مندی چون چشم او ناشی می‌شود در حالی که روبات دارای چنان سنسوری نیست. از آن‌جا که طراح در طول زندگی‌اش با سنسور بینایی قوی‌ای چون چشم کار کرده است و همچنین دینامیک بدن او کاملاً متفاوت با دینامیک روبات است، به هنگام طراحی چنان رفتاری برای روبات می‌بایست حدس بزند که چگونه با استفاده از سونار می‌توان آن رفتار را ایجاد کرد. این مشابه‌سازی ذهنی به دور از هرگونه بهینه‌گی‌ای است و در به‌ترین حالت او تنها به طراحی قابل قبول‌ای دست می‌یابد.

مساله‌ی یادگیری رفتار برای حل این مشکل مطرح می‌شود. هدف یادگیری رفتار پیدا کردن نگاشت مناسب بین ورودی هر رفتار و خروجی آن است. مناسب بودن یک رفتار به تاثیر آن بر

عمل کرد عامل باز می‌گردد. به بیان‌ای دیگر، هدف پیدا کردن مجموعه پارامترهای توصیف‌گر رفتاری است که مساله‌ی بهینه‌سازی (۲-۹) را حل کند.

هدف این فصل معرفی‌ی روش‌ای برای انجام این کار است. در این فصل فرض می‌کنم که برای حل مساله‌ی یافتن نگاشت $a_i = B_i(s'_i)$ مطلوب، فضای حالت (S_i) و عمل (A_i) هر رفتار و هم‌چنین سیگنال تقویت دانسته شده است.^۱

در ابتدا نشان خواهیم داد که چگونه می‌توان ارزش کل عامل را به اجزای هر رفتار تقسیم کرد. این قدم شبیه به آن چیزی است که در یادگیری ساختار انجام دادم. بعد از این مرحله‌ی اولیه، چارچوب‌ای کلی برای حل مساله‌ی تقسیم امتیاز در خانواده‌ای از سیستم‌های چندعامله معرفی می‌کنم و سپس با استفاده از آن چگونگی تقسیم امتیاز بین رفتارها در معماری‌ی سلسله‌مراتبی‌ی PPSSA را مشخص خواهیم کرد. در نهایت بررسی‌ای بر تاثیر وجود NA بر بازده عامل انجام خواهیم داد و روابطی برای یادگیری و به روز رسانی‌ی دانش هر رفتار نسبت به محیط اطراف‌اش ارائه خواهیم کرد.

۲-۴ تجزیه ارزش به اجزای رفتاری

در این بخش به ارتباط بین تجربه‌ی هر رفتار که به صورت تابع ارزش حالت-عمل $(Q_i(s'_i, a_i))$ ثبت می‌شود و ارزش کل عامل می‌پردازم. نشان می‌دهم که ارزش کل معماری به صورت مجموعه‌ای از ارزش‌های اجزای هر رفتار (یعنی حالت‌های هر رفتار) قابل بیان است.

در (۳-۶)، V_T به اجزای ساده‌تری که شامل V_{ij} ، $P\{B_j|L_i\}$ و $P(L_i \text{ is controlling})$ بودند تجزیه شد. اگر ارزش هر حالت هر رفتار را به صورت امید ریاضی پاداش لحظه‌ای آن حالت در نظر بگیریم (در نتیجه $V_j(s'_j)$ یا $Q_j(s'_j, a_j)$ تخمین‌ای از متوسط پاداش در طول یک اپیزود یا پاداش discount شده نیست)، آن‌گاه V_{ij} را به صورت مجموعه‌ای از اجزای $V_j(s'_j)$ می‌توان نوشت. این محدودیت به دلیل تعریف V_{ij} (۳-۴) است که در آن V_{ij} متوسطی از پاداش‌هایی است که B_i ‌ی که درآمین لایه بوده است هنگام کنترل‌کننده‌بودگی‌اش دریافت کرده است. حالت‌های کلی‌تر توابع ارزش حالت و حالت-عمل در بخش‌های بعدی استفاده می‌شوند اما این تجزیه تنها هنگامی که ارزش را به

^۱ نتایج این فصل و نتایج تکمیل شده‌ی فصل یادگیری ساختار در [Farahmand05b] آمده است.

صورت پاداش لحظه‌ای تعریف می‌کنیم معتبر است. با این فرض، V_{ij} بازنمایی مرتبه صفر به صورت زیر قابل گسترش است^۲:

$$V_{ij} = \sum_{s'_j \in S'_j} P\{s'_j | B_j \text{ is controlling in } L_i\} V_j(s'_j) \quad (۴-۱)$$

که در آن فضای ورودی تحریک‌شده‌ای است که توسط M_j از فضای اصلی $S_j \subset S$ نگاشت شده است. همچنین داریم:

$$V_j(s'_j) = \sum_{a_j \in A_j} \pi_j(s'_j, a_j) Q_j(s'_j, a_j) \xrightarrow{\text{greedy policy}} = \max_{a_j} Q_j(s'_j, a_j) \quad (۴-۲)$$

که در آن $\pi_j(s'_j, a_j)$ احتمال انتخاب عمل a_j در حالت s'_j رفتار B_j است در آن هنگام که آن رفتار کنترل‌کننده باشد. با توجه به این تجزیه و در نظر گرفتن تجزیه‌ی رفتاری‌ای که در بازنمایی مرتبه صفر داشتیم، V_T را می‌توان بر حسب $P\{B_j | L_i\}$ ، $P(L_i \text{ is controlling})$ ، $P\{s | B_j \text{ is controlling in } L_i\}$ و $Q_j(s', a_j)$ نوشت^۳:

$$V_T = \sum_{i=1}^m \sum_{j=1}^n \sum_{s' \in S'_j} \sum_{a_j \in A_j} P\{B_j | L_i\} P(L_i \text{ is controlling}) P\{s' | B_j \text{ is controlling in } L_i\} \pi_j(s'_j, a_j) Q_j(s'_j, a_j) \quad (۴-۳)$$

به دست آوردن رابطه‌ای مشابه برای بازنمایی مرتبه یک نیز قابل انجام است.

^۲ اگر S'_j فضایی پیوسته باشد، جمع کردن به انتگرال‌گیری تبدیل می‌شود. مثال‌ای از استفاده از انتگرال به جای جمع در بخش "بهبه‌نگی"،

مجموعه عمل‌ها و به روز رسانی ارزش^۳ این فصل آمده است.

^۳ توجه کنید که همه‌ی این مقادیر را می‌توان در طول یادگیری تخمین زد.

۴-۳) تخصیص اعتبار به رفتارها

یکی از مسایل اساسی در یادگیری رفتار، چگونگی تخصیص اعتبار به رفتارها است. به طور معمول معماری رفتارگرا شامل چندین رفتار مختلف است که هر کدام نقش‌ای برای ایجاد رفتار مطلوب کلی ایفا می‌کنند. اگر هر کدام از رفتارها را یک عامل بدانیم، یک معماری رفتارگرا را می‌توان به صورت سیستمی چندعامله در نظر گرفت و مساله تخصیص اعتبار در معماری رفتارگرا را به صورت تخصیص اعتبار در سیستمی چندعامله در نظر گرفت.

با این وجود مساله تخصیص اعتبار در سیستم‌های چندعامله نیز چندان ساده نیست چون در حالت کلی اطلاعات کافی‌ای برای چگونگی تقسیم اعتبار تیمی به هر کدام از عامل‌ها در اختیار نیست. به عبارت دیگر در حالت کلی مشخص نیست که نقش هر کدام از عامل‌ها در تنبیه یا پاداش دریافتی چه بوده است ([Harati04]). همچنین در [Shoham04] نگاه نقادانه‌ای به روش‌های مختلف یادگیری چندعامله تقویتی شده است. با این وجود گاهی اوقات قیدهایی که در ساختار تیم چندعامله وجود دارد ما را در تقسیم امتیاز راهنمایی می‌کند. در این بخش ابتدا چارچوبی را برای تقسیم امتیاز در خانواده‌ای محدود از سیستم‌های چندعامله معرفی می‌کنم و سپس با استفاده از آن چارچوب، روشی برای حل مساله تقسیم امتیاز در PPSSA پیش‌نهاد خواهم کرد.

۴-۳-۱) سیستم چندعامله به عنوان یک جبر بول^۴

فرض کنید که عامل‌های $A_1, A_2, \dots$ و A_n با هم دیگر تیم چندعامله‌ی T را می‌سازند. همچنین فرض کنید که موفقیت یا شکست تیم T تابعی از موفقیت یا شکست هر کدام از A_i باشد. به عنوان مثال $T = A_1 \wedge A_2$ تیمی است که شرط موفقیت آن، صحیح عمل کردن A_1 و A_2 باشد (حالت و) و $T = A_1 \vee A_2$ تیمی است که شرط موفقیت آن صحیح کار کردن حداقل یکی از A_1 یا A_2 باشد (حالت یا). تیم‌های "و"ی و "یا"ی دو حد نهایت مختلف سیستم‌های چندعامله‌ی هم‌کار^۵ هستند.

^۴ Boolean Algebra

^۵ Cooperative Multi-Agent Systems

روش‌هایی برای حل این دو حالت حدی در [Harati04] پیشنهاد شده است. آن روش‌ها پیشنهاد می‌کنند که هنگامی که تیم "و" ای ("یا" ای) پاداش (تنبیه) دریافت می‌کند، آن پاداش (تنبیه) می‌بایست به همه‌ی عامل‌ها تخصیص داده شود. روش‌های تقریبی‌ای برای تخصیص اعتبار در زمان‌هایی که تیم "و" ای ("یا" ای) تنبیه (پاداش) دریافت می‌کند نیز پیشنهاد شده است که مبتنی بر استفاده از معیارهای خبرگی^۶ ([Nili02]) هر عامل است. در این پایان‌نامه تنها از نتایج "و" ای-پاداش و "یا" ای-تنبیه استفاده کرده‌ام.

در مسایل‌ای که موفقیت یا شکست تیم T تابع‌ای منطقی از موفقیت یا شکست هر کدام از اجزای‌ای (A_i) باشد، تخصیص اعتبار از قوانین جبر بول تبعیت می‌کند و در نتیجه همه‌ی عملیات تعریف شده و استنتاج‌های ممکن در آن جبر قابل استفاده خواهد بود. به عنوان مثال دانستن این که تیم $T = A_1 \wedge A_2$ درست کار کرده است ما را قادر به این می‌کند که بدانیم هر کدام از A_1 و A_2 نیز به تنهایی درست کار کرده‌اند.

مجموعه‌ی B را که شامل همه‌ی ترکیب‌های عامل‌های $\{A_i\}$ است در نظر بگیرید. اصول جبر بول B برای تمام سیستم‌های چندعامله‌ای که صحت کارکرد تیم با عبارت‌ای منطقی به صحت کارکرد هر کدام از اجزای‌اش ارتباط دارد صادق است:

ب^۰: اگر A_1 و A_2 در مجموعه‌ی B باشند، $A_1 \wedge A_2$ و $A_1 \vee A_2$ نیز در مجموعه‌ی B هستند. این خاصیت از تعریف B مشخص است.

ب^۱ - تعویض‌پذیری^۷: عمل کرد T مستقل از ترتیب در نظر گرفتن هر کدام از عامل‌های آن است، یعنی:

^۶ Expertness
^۷ Commutativity

$$A_1 \wedge A_2 \equiv A_2 \wedge A_1 \quad (4-4)$$

$$A_1 \vee A_2 \equiv A_2 \vee A_1 \quad (4-5)$$

ب ۲ - شرکت‌پذیری^۸: عمل کرد T مستقل از مکان شروع محاسبه‌ی عمل کرد هر کدام از عامل‌های آن است. به عبارت دیگر اهمیت‌ای ندارد که اول درستی‌ی A_1 و A_2 بررسی شود و سپس نتیجه به درستی‌ی A_3 "و" شود و یا این که ابتدا عمل کرد A_2 و A_3 بررسی شود و سپس نتیجه به درستی‌ی A_1 "و" شود.

$$A_1 \wedge (A_2 \wedge A_3) \equiv (A_1 \wedge A_2) \wedge A_3 \quad (4-6)$$

$$A_1 \vee (A_2 \vee A_3) \equiv (A_1 \vee A_2) \vee A_3 \quad (4-7)$$

ب ۳ - توزیع‌پذیری^۹: "و" ($\wedge$) و "یا" ($\vee$) به طور متقابل قابل توزیع هستند:

$$A_1 \wedge (A_2 \vee A_3) = (A_1 \wedge A_2) \vee (A_1 \wedge A_3) \quad (4-8)$$

$$A_1 \vee (A_2 \wedge A_3) = (A_1 \vee A_2) \wedge (A_1 \vee A_3) \quad (4-9)$$

سیستم چندعامله‌ای را مطابق با سمت چپ (۴-۸) در نظر بگیرید. اگر A_1 درست کار کند و دست‌کم یکی از A_2 یا A_3 درست کار کنند، آن‌گاه کل سیستم هم درست کار می‌کند. فرض کنید A_2 عامل‌ای است که درست کار می‌کند. در نتیجه درست کار کردن $A_1 \wedge A_2$ قابل استنتاج است. از این نتیجه می‌شود که $(A_1 \wedge A_2) \vee (A_1 \wedge A_3)$ نیز درست کار می‌کند و در نتیجه تساوی دو طرف رابطه‌ی (۸-۴) برقرار خواهد شد. حالت‌های دیگر (درست یا نادرست بودن هر کدام از A_1 ، A_2 و A_3) نیز به همین ترتیب قابل استنتاج است. اثبات صحت (۴-۹) نیز به روش‌ای مشابه انجام می‌شود.

⁸ Associativity

⁹ Distributivity

ب۴ - عامل‌های خنثی^{۱۰}: عامل‌های خنثی خیالی 1 و 0 ای برای اپراتور "و" ($\wedge$) و "یا" ($\vee$) وجود دارند. 1 عامل مجازی‌ای است که در همه‌ی حالات درست کار می‌کند و 0 نیز عامل مجازی‌ی دیگری است که در هیچ شرایطی درست کار نمی‌کند. روابط زیر برای این عامل‌ها برقرار است:

$$A \wedge 1 = A \quad (۴-۱۰)$$

$$A \vee 0 = A \quad (۴-۱۱)$$

ب۵ - عامل‌های متمم^{۱۱}: برای هر عامل اپراتور A ، عامل مجازی $\bar{A}$ وجود دارد که هر هنگام این دو با هم "یا" ($\vee$) شوند، نتیجه همیشه درست خواهد بود. برای عامل A ای که درست عمل می‌کند، عامل متمم $\bar{A}$ آن نادرست عمل می‌کند و برای عامل‌ای که نادرست عمل می‌کند، عامل متمم $\bar{A}$ درست عمل می‌کند. عامل متمم $\bar{A}$ برای حالت "و" ای عامل‌ای است که $A \wedge \bar{A}$ را نادرست می‌کند.

۴-۳-۲) تخصیص اعتبار به رفتارها در معماری Subsumption

حل مسالهی تخصیص اعتبار به رفتارهای معماری رفتارگرایی PPSSA به کمک چارچوب معرفی شده ممکن است. همان‌طور که پیش‌تر گفتیم معماری رفتارگرا را می‌توان به صورت سیستم‌ای چندعامله در نظر گرفت. نشان خواهیم داد که درستی یا نادرستی عمل‌کرد معماری مورد نظر تابع‌ای منطقی از عمل‌کرد هر کدام از رفتارها است و در نتیجه استفاده از چارچوب معرفی شده ممکن خواهد بود.

فرض کنید B^* بیان‌گر رفتار کنترل‌کننده در لحظه‌ی مورد بررسی باشد، $\{B_{u_i}(NA)\}$ مجموعه‌ی رفتارهای تحریک‌شده‌ی بالاتر از B^* در معماری T که خروجی NA داده‌اند و $\{B_{\bar{u}_i}\}$ نیز رفتارهای بالاتر از B^* ای که تحریک نشده‌اند. همچنین $\{B_{l_i}\}$ آن رفتارهای پایین‌تر از B^* ای هستند که فعال شده‌اند (یعنی خروجی‌ای غیر از NA پیش‌نهاد داده‌اند) اما توسط B^* سرکوب شده‌اند (در نتیجه

¹⁰ Neutral

¹¹ Complement

یاری‌بهری^{۱۲} در خروجی نداشته‌اند) و $\{B_{\bar{l}_i}\}$ نیز رفتارهای پایین‌تری هستند که اصلاً تحریک نشده‌اند. مشخص است که این مجموعه همه‌ی رفتارها را افراز^{۱۳} می‌کند - $\{B^*, \{B_{u_i}(NA)\}, \{B_{\bar{u}_i}\}, \{B_{l_i}\}, \{B_{\bar{l}_i}\}\}$ همه‌ی رفتارها را یک و فقط یک بار شامل می‌شود. برای حل مسالهی تخصیص اعتبار در این معماری دو حالت مختلف جداگانه تحلیل می‌شوند: زمان‌هایی که عامل پاداش دریافت کرده است و زمان‌هایی که عامل تنبیه شده.

فرض کنید که معماری T در لحظه‌ی جاری پاداش‌ای دریافت کند. این پاداش حاصل عمل صحیح B^* و فعال نشدن $\{B_{u_i}(NA)\}$ بوده است. در نتیجه اعضای مجموعه‌ی $\{B^*, \{B_{u_i}(NA)\}\}$ به درستی عمل کرده‌اند. پس عبارت جبر بول زیر در حالت دریافت پاداش صحیح است:

$$T_{R^+} : B^* \wedge (B_{u_1}(NA) \wedge \dots \wedge B_{u_k}(NA)). \quad (۴-۱۲)$$

همچنین اعضای مجموعه‌ی $\{B_{\bar{u}_i}\}$ یاری‌بهری در بازده کلی سیستم نداشته‌اند؛ آن‌ها هیچ تصمیم‌گیری‌ای انجام نداده‌اند که خوب یا بد باشد. پس می‌توان آن‌ها را به عنوان عامل‌های خنثی در نظر گرفت. عامل خنثی‌ی اپراتور "و"، با 1 نمایش داده می‌شود. پس می‌توان نوشت:

$$T_{R^+} : B^* \wedge (B_{u_1}(NA) \wedge \dots \wedge B_{u_k}(NA) \wedge 1). \quad (۴-۱۳)$$

در مورد رفتارهای پایین $\{B_{l_i}\}$ و $\{B_{\bar{l}_i}\}$ می‌توان گفت که هیچ تأثیری در عمل کرد سیستم نداشته‌اند و در نتیجه حرف خاص‌ای در مورد درست یا نادرست عمل کردن آن‌ها نمی‌توان گفت:

$$T_{R^+} : \underbrace{\hat{1}}_{\text{lower behaviors}} \wedge \underbrace{\bar{B}^*}_{\text{controlling behavior}} \wedge \overbrace{\left(B_{u_1}(NA) \wedge \dots \wedge B_{u_k}(NA) \wedge \underbrace{\hat{1}}_{\text{not excited}} \right)}^{\text{upper behaviors}}. \quad (۴-۱۴)$$

¹² Contribution

¹³ Partition

اینک تخصیص اعتبار به رفتارها در حالت سیگنال تقویت مثبت مشخص شده است: در این حالت $T_{R^+} \equiv 1$ و در نتیجه استنتاج روبرو با توجه به خواص جبر بول صحیح است^{۱۴}:

$$T_{R^+} : + \Rightarrow \begin{cases} B^* : + \\ \{B_{u_i}(NA)\} : + \\ \{B_{\bar{u}}\} : \text{unknown} . \\ \{B_l\} : \text{unknown} \\ \{B_{\bar{l}}\} : \text{unknown} \end{cases} \quad (۴-۱۵)$$

راه حل مسالهای تخصیص اعتبار در حالت دریافت سیگنال تنبیه اندکی گنگ و نامشخص است. معماری $T = [B_{l_i} \quad B_{\bar{l}_i} \quad B^* \quad B_{u_i}(NA) \quad B_{\bar{u}_i}]$ را در نظر بگیرید که سیگنال تنبیه دریافت کرده است. در حالت سیگنال مثبت، درستی عمل کرد معماری را بدین گونه به درستی عمل کرد اجزای اش ارتباط داده بودم:

$$(T \equiv +) \Rightarrow [B^* \wedge (B_{u_1}(NA) \wedge \dots \wedge B_{u_k}(NA))] \equiv + . \quad (۴-۱۶)$$

اما از نظر منطقی نمی توان عبارت زیر را به طور مستقیم با عوض کردن علامت مثبت به منفی نتیجه گرفت:

$$(T \equiv -) \Rightarrow [B^* \wedge (B_{u_1}(NA) \wedge \dots \wedge B_{u_k}(NA))] \equiv - . \quad (۴-۱۷)$$

^{۱۴} قابل ذکر است که این روش تخصیص اعتبار ممکن است به پاسخهای زیربهنه ای منجر شود. دلیل چنین چیزی این است که ممکن است اگر رفتارهای بالاتر $B_{u_i}(NA)$ به جای NA عمل ای پیش نهاد می دادند نتیجه ی به تری از عمل پیش نهادی B^* حاصل می شد. این مشکل در مسایل ای که در آن بیش از دو سطح درست و نادرست (یا مثبت و منفی) عمل کرد سیستم وجود دارد رخ می دهد. به نظر می رسد اگر سیستم جستجویی کاوش گرانه در انتخاب عمل های اش داشته باشد این مساله رفع شده و سیستم قادر به یافتن پاسخ های بهینه نیز شود.

برای تحلیل حالت تنبیهی، فرض کنید که عمل صحیح که منجر به پاداش می‌شود - که البته از پیش دانسته نیست - در فضای عمل همه‌ی رفتارهای B^* و $\{B_{u_i}(NA)\}$ وجود داشته باشد. یعنی:

$$a^* \in \begin{cases} A_{B^*} \\ A_{B_{u_i}(NA)} \end{cases} \quad \forall B_{u_i}(NA) \quad (4-18)$$

در نتیجه هنگام دریافت تنبیه، ممکن است همه‌ی رفتارهای $\{B_{u_i}(NA)\}$ با پیش‌نهاد NA شان درست عمل کرده باشند اما ایراد از عمل رفتار B^* بوده باشد. اگر از $\bar{B}$ به عنوان نماد رفتاری که پاسخ نادرستی داده است بهره جوییم، این وضعیت به صورت $\bar{B}^* B_{u_1} \dots B_{u_k}$ قابل بیان است. حالت دیگر می‌تواند این باشد که $B_{u_1}(NA)$ به نادرستی خروجی داده (یعنی انتخاب NA انتخاب درستی نبوده است) و به جای‌اش باید $a_1 = a^* \in A_{B_{u_1}}$ را انتخاب می‌کرده است. اگر نماد $\hat{B}$ را به عنوان نماد رفتار بدون یاری‌بهر در نتیجه‌ی حاصل در لحظه‌ی فعلی در نظر بگیریم، این شرایط به صورت $\bar{B}^* \bar{B}_{u_1} \dots B_{u_k}$ قابل بیان است. در حالت کلی یکی از شرایط زیر رخ می‌دهد:

$$\begin{cases} \bar{B}^* B_{u_1} \dots B_{u_k} \\ \hat{B}^* \hat{B}_{u_1} \hat{B}_{u_2} \dots \bar{B}_{u_i} \dots B_{u_k} \end{cases} \quad i = 1, \dots, k \quad (4-19)$$

پس با فرض انجام شده همه‌ی رفتارهای $\{B^*, \{B_{u_i}(NA)\}\}$ می‌توانستند به گونه‌ای عمل کنند که سیستم پاداش دریافت کند. چون در هنگام دریافت تنبیه هیچ‌کدام از آن‌ها این‌گونه عمل نکرده‌اند، پس می‌توان نتیجه گرفت که همه‌شان می‌بایست تنبیه شوند. به مانند قبل، چیزی در مورد رفتارهای بی‌یاری‌بهر $\{B_{\bar{u}_i}\}$ ، $\{B_{l_i}\}$ و $\{B_{\bar{l}_i}\}$ نمی‌توان گفت: پس آن‌ها به مانند عامل‌های خنثی در سیستم چندعامله عمل می‌کنند. روی هم رفته می‌توان نوشت:

$$T_{R^-} : \underbrace{\vec{0}}_{\text{lower behaviors}} \vee \underbrace{\vec{B}^*}_{\text{controlling}} \vee \underbrace{\left(B_{u_1}(NA) \vee B_{u_2}(NA) \vee \dots \vee B_{u_k}(NA) \vee \vec{0} \right)}_{\substack{\text{upper behaviors} \\ \text{not excited}}} \quad (4-20)$$

در نتیجه تخصیص اعتبار در حالت دریافت سیگنال تنبیه در صورتی که a^* عضوی از فضای عمل همه‌ی رفتارهای B^* و $\{B_{u_i}(NA)\}$ باشد بدین گونه خواهد بود:

$$T_{R^-} : - \Rightarrow \begin{cases} B^* : - \\ \{B_{u_i}(NA)\} : - \\ \{B_u\} : \text{unknown} . \\ \{B_l\} : \text{unknown} \\ \{B_f\} : \text{unknown} \end{cases} \quad (4-21)$$

نتیجه در ظاهر و با تغییر علامت مثبت به منفی شبیه به حالت سیگنال تقویت مثبت است اما فرآیند استنتاج پشت‌اش اندکی متفاوت است. در حالت منفی فرض کرده بودم که همه‌ی رفتارها، عمل صحیح را در فضای عمل‌های‌شان دارند (4-18) ولی لزومی به چنین فرض‌ای در حالت مثبت نیست. گاهی اوقات این فرض انجام شده چندان صحیح نیست و ممکن است بعضی از رفتارهای تحریک شده‌ی بالایی $\{B_{u_i}(NA)\}$ قادر به خروجی صحیح دادن نباشند (یعنی $a^* \notin A_{B_{u_i}}$). در نتیجه تنبیه کردن آن‌ها عاقلانه نخواهد بود و به عنوان نویزی در سیگنال تقویت محسوب می‌شود.¹⁵ به هر حال از پیش دانسته نیست که آیا عمل بهینه در مجموعه‌ی اعمال یک رفتار خاص وجود دارد یا نه و در نتیجه اشتباه در پاداش نادرست اجتناب‌ناپذیر است. البته در مسایل‌ای که خروجی همه‌ی رفتارها یک‌سان است – مثلاً همه به کنترل سرعت چرخ‌های روبات مشغول‌اند – و یا این که فضای تحریک رفتارها جدا از هم است – و در نتیجه در هر زمان تنها یک رفتار تحریک می‌شود – چنین مشکل‌ای پیش نمی‌آید.

¹⁵ برای مشاهده‌ی تحلیل‌هایی درباره‌ی نتایج تأثیر عدم قطعیت سیگنال تقویت و وجود نویز در آن بر تابع ارزش و سیاست عامل به [Singh94]، [Williams94] و [Farahmand05a] مراجعه کنید. نتایج مرجع آخر – به عنوان بخشی‌ای از کار انجام شده در این پژوهش – در فصل ۸ آورده می‌شود.

به طور خلاصه از تحلیل‌های این بخش به این نتیجه می‌رسیم که برای حل مسأله‌ی تخصیص اعتبار در PPSSA، عمل پیش‌نهادی رفتار کنترل‌کننده‌ی B^* و همه‌ی NA های رفتارهای بالاتر $\{B_{u_i}(NA)\}$ می‌بایست متناسب با سیگنال تقویت دریافتی عامل تقویت شوند.

۴-۴) بهینه‌گی، مجموعه عمل‌ها و به روز رسانی ارزش

در این بخش به بررسی‌ی تاثیر وجود NA بر بهینه‌گی عمل‌کرد عامل می‌پردازیم و نشان خواهیم داد که وجود آن عمل مجازی شانس بالا رفتن عمل‌کرد عامل را ایجاد می‌کند. پس از آن روش‌ای برای به روز رسانی ارزش‌های هر حالت رفتار برای دو وضعیت مختلف پاداش لحظه‌ای و پاداش بازگشته معرفی خواهیم کرد.

فرض کنید می‌خواهیم ترتیب بهینه‌ی دو رفتار B_1 و B_2 را -که به ترتیب دارای زیرفضای تحریک S_1 و S_2 هستند- بیابیم (مسأله‌ی یادگیری ساختار). مشخص است اگر عمل NA در مجموعه‌ی اعمال انتخابی آن دو رفتار نباشد، آن‌ای که بالاتر است (به آن B_{higher} می‌گوییم که می‌تواند هر کدام از آن دو رفتار باشد) کنترل کامل عامل را در همه‌ی حالت‌هایی از فضای اصلی ($s \in S$) که جزو زیرفضای تحریک‌اش است می‌گیرد ($\{s | \forall s \in S; s \in S_{higher}\}$). رفتار دیگر (B_{lower}) در آن مجموعه سرکوب می‌شود و هیچ کنترلی در اختیار ندارد. برای این‌که بفهمیم این ترتیب خاص بهینه است یا نه، یاری‌بهر هر کدام از رفتارها بر پاداش دریافتی را با هم مقایسه می‌کنیم. منظور از یاری‌بهر یک رفتار در پاداش دریافتی عامل، پاداش‌ای است که عامل به دلیل سیاست خاص آن رفتار به دست آورده است. این پاداش ممکن است به دلیل عمل آن رفتار و دریافت پاداش لحظه‌ای باشد و یا این‌که به دلیل فعال نشدن و "اجازه‌ی کنترل‌کننده‌بودن به رفتارهای پایین دادن" باشد. یعنی یک رفتار اگر خود نیز کنترل‌کننده نباشد ولی با انتخاب NA در زمان مناسب به رفتار پایینی‌اش اجازه‌ی انجام عمل مناسب بدهد قابل تقدیر تصمیم گرفته است. همچنین این پاداش ممکن است تاخیر یافته یا لحظه‌ای باشد. برای روشن شدن معنای یاری‌بهر، فرض کنید B_{higher} کنترل عامل را در اختیار گرفته باشد و عامل پاداش دریافت کند. مشخص است که این پاداش به دلیل یاری‌بهر B_{higher} بوده است. حال فرض کنید که B_{higher} تحریک می‌شود ولی NA پیش‌نهاد می‌کند و در نتیجه پس از آن B_{lower}

کنترل کننده می شود. اگر عامل پاداش دریافت کند —همان طور که در زیربخش پیشین گفتیم— نه فقط به دلیل عمل کرد مناسب عمل انتخابی B_{lower} بوده که به دلیل پیش نهاد مناسب NA B_{higher} نیز بوده است. "پاداش ای که به دلیل یاری بهر B_i در معماری T در زیرفضای اکید^{۱۶} $S^* (\subseteq S)$ حاصل شده است" را بدین گونه تعریف می کنیم:

$$\begin{aligned} \bar{r}_i|_{s \in S^*} &= E[R_i(s)]|_{s \in S^*} = E[R(s) | \text{Reward is achieved by the contribution of } B_i \wedge s \in S^*] \\ &= \int_{s \in S^*} R(s) p\{B_i \text{ is excited in } s\} p\{s \in S^*\} ds \end{aligned} \quad (۴-۲۲)$$

که در آن $p\{s \in S^*\}$ احتمال قرارگیری حالت عامل در s است به شرط آن که حالت عامل در مجموعه S^* باشد. $R(s)$ نیز پاداش بازگشته عامل از حالت s است. پاداش بازگشته می تواند پاداش تک گام^{۱۷} (که ارزش حالت های بعدی در ارزش حالت فعلی تاثیری ندارد) و یا پاداش بازگشته ی چندگام ای^{۱۸} مانند متوسط پاداش های اپیزود باشد.

اینک با تعاریف انجام شده می توان دو رفتار را در کل فضای حالت عامل (S) بدین گونه مقایسه کرد:

$$\begin{aligned} \int_S R(s) p\{B_{higher} \text{ is excited in } s\} p\{s \in S\} ds &> \int_S R(s) p\{B_{lower} \text{ is excited in } s\} p\{s \in S\} ds \\ \Leftrightarrow B_{higher} &> B_{lower} \end{aligned} \quad (۴-۲۳)$$

این مقایسه مشخص می کند که کدام رفتار یاری بهر بیش تری در پاداش دریافتی عامل در طول زندگی اش داشته است. اما با این وجود این ترکیب ممکن است به ازای زیرمجموعه ی

^{۱۶} زیرفضای اکید S^* (strict subspace) بدین گونه تعریف می شود: $S^* = \{s | \forall s \in S; M_i(s) \in S'^*\}$

^{۱۷} Single Step

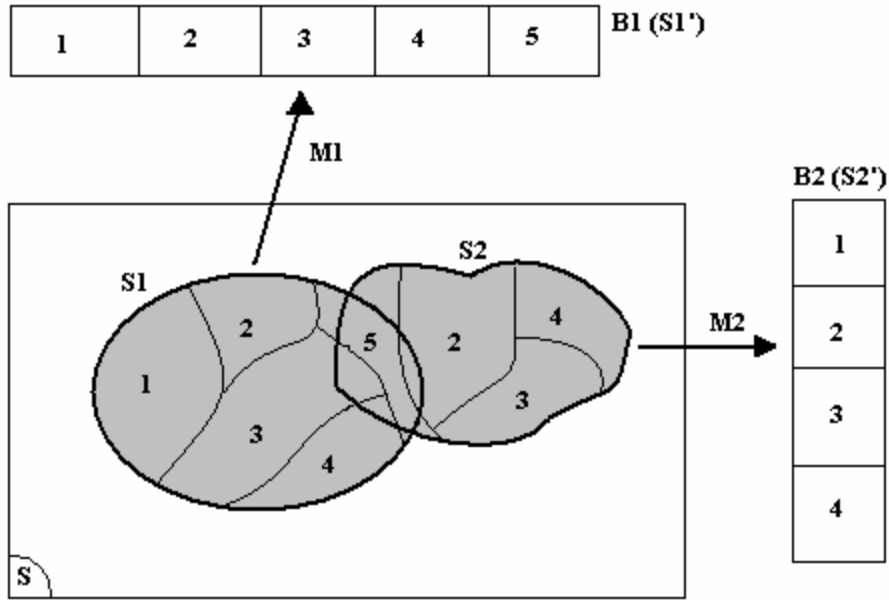
^{۱۸} Multi Step

$S^* \subset S_{higher} \cap S_{lower}$ ای که $\bar{r}_{lower}|_{s \in S^*} > \bar{r}_{higher}|_{s \in S^*}$ بهینه نباشد. اگر رفتار B_{higher} همیشه رفتار B_{lower} را بدون توجه به برتری دومی در بعضی از زیرمجموعه‌های فضای حالت سرکوب کند، بازده عامل کاهش می‌یابد. این مشکل با تعریف NA قابل حل است: عمل مجازی‌ای که در صورت انتخاب کاری انجام نمی‌دهد و اجازه می‌دهد که رفتار با وجود تحریک شدن، فعال نشود و شانس کنترل را به رفتارهای پایین واگذارد. این دلیل وجودی NA و تعمیم فضای عمل رفتارها است.

برای مشخص شدن این که در چه شرایطی B_{higher} می‌بایست کنترل کننده باشد و در چه شرایطی اجازه‌ی کنترل کننده‌گی به رفتارهای پایین بدهد، یاری‌بهر آن رفتار می‌بایست با یاری‌بهر رفتارهای پایین‌ترش مقایسه شود. برای این مقایسه، یاری‌بهر عمل‌ای از آن رفتار با یاری‌بهر انتخاب NA همان رفتار - که معادل یاری‌بهر رفتارهای پایین‌تر است - مقایسه می‌شود. در این مقایسه می‌بایست به این نکته توجه کرد که فضای تحریک رفتارهای متفاوت یک‌سان نیست ($S_i \neq S_j; i \neq j$) و همچنین از آن‌جا که در حالت کلی $M_i \neq M_j$ ، گاهی اوقات B_i و B_j ممکن است با هم تحریک شوند و گاهی تنها یکی از آن‌ها (شکل ۴-۱). به بیان‌ای دیگر، ممکن است به ازای زیرفضاهایی ($S_1, S_2 \subset S$) از فضای حالت $s'_i = M_i(S_1)$ و $s'_j = M_j(S_1)$ (پس B_i و B_j هر دو تحریک شده‌اند) ولی $s'_i = M_i(S_2)$ اما $s'_j \neq M_j(S_2)$ (یعنی B_j دیگر در آن حالت درونی‌ی s'_j تحریک شده نیست). در نتیجه در معماری‌ی T ای با $B_k > \dots > B_i$ ، مقایسه‌ی شایستگی‌ی B_i و رفتارهای پایین‌ترش می‌بایست در فضای تحریک مشترک آن رفتارها باشد. این نکته را در تحلیل‌های بعدی در نظر می‌گیرم.

رفتار B_i را در نظر بگیرید. $S^\circ(s'_i) \subset S$ را به صورت زیرمجموعه‌ای از S در نظر بگیرید که به یک نقطه s'_i از بازنمایی داخلی رفتار B_i نگاشت می‌یابد (هر کدام از بخش‌های شکل ۴-۱). تعریف $S^\circ(s'_i)$ بدین گونه است:

$$S^\circ(s'_i) = \{s | \forall s \in S; M_i(s) = s'_i\} \quad (4-24)$$



شکل ۴-۱. زیرفضاهای تحریک‌پذیر رفتارهای B_1 و B_2 در فضای S و تناظر آن‌ها با S'_1 و S'_2 . توجه کنید که زیرفضاهای تحریک‌پذیر ممکن است هم‌پوشانی داشته باشند.

حال می‌توان $\bar{r}_i|_{s \in S^\circ}$ را برای $S^\circ(s'_i)$ بدین‌گونه نوشت:

$$\bar{r}_i|_{s \in S^\circ} = \int_{s \in S^\circ} R(s) p\{B_i \text{ is excited in } s\} p\{s \in S^\circ\} ds = R(s') p\{B_i \text{ is excited in } s'\} = V_i(s'_i) \quad (۴-۲۵)$$

متوسط یاری‌بهر رفتار B_i در شرایطی که در حالت s'_i بوده است $V_i(s'_i)$ نام‌گذاری می‌شود. به نحوی مشابه می‌توان $Q_i(s', a_i)$ را به صورت امید یاری‌بهر رفتار B_i در حالت درونی s' پس از انتخاب عمل a_i بدین‌صورت نوشت:

$$Q_i(s', a_i) = \int_{s \in S^\circ} R(s) p\{B_i \text{ is excited in } s | B_i \text{ selects } a_i\} p\{s \in S^\circ\} ds = R(s') p\{B_i \text{ is excited in } s' | B_i \text{ selects } a_i\} \quad (۴-۲۶)$$

و ارتباط $Q_i(s', a_i)$ با $V_i(s')$ نیز بدین‌گونه خواهد بود:

$$V_i(s') = E_{a_i \sim \pi_i(s')} [Q_i(s', a_i)]. \quad (۴-۲۷)$$

برای این که مشخص شود آیا رفتارهای پایینی در حالت ای مانند s'_i می توانند به تر عمل کنند یا این که به تر است خود آن رفتار عمل ای پیش نهاد کند، می بایست بازده رفتارهای پایینی را در حالت هایی از فضای اصلی که B_i را در s'_i قرار داده (و در نتیجه B_i تحریک شده است) و دست کم یکی از آنها (رفتارهای پایینی) تحریک شده باشد با بازده انتخاب یک عمل از B_i مقایسه کرد. این کار با مقایسه ی $Q_i(s', a_i)$ و $Q_i(s', NA)$ - که امید یاری بهر انتخاب NA و امید یاری بهر رفتارهای پایینی است - انجام می شود. $Q_i(s', NA)$ مشابه با $Q_i(s', a_i)$ بدین گونه مشخص می شود:

$$Q_i(s', NA) = \int_{s \in S^\circ} R(s) p\{B_i \text{ is excited in } s | B_i \text{ selects } NA\} p\{s | s \in S^\circ\} ds. \quad (۴-۲۸)$$

توجه کنید که به هنگام انتخاب NA ، B_i کنترل را به رفتارهای پایینی واگذار می کند. در نتیجه یاری بهر آن رفتار ناشی از این کار است و نه عمل ای که خود مستقیم به عامل ارسال کرده است.

برای بهینه عمل کردن می بایست ارزش انتخاب NA را با انتخاب عمل ای از مجموعه اعمال A_i مقایسه کرد: اگر ارزش عمل نکردن بالاتر بود، با شانس کنترل به رفتارهایی پایینی دادن به بازده بالاتری - از دیدی آماری - خواهیم رسید. انتخاب بین اعضای A'_i (که شامل NA نیز می شود) به مانند معمول است و می توان از روش های انتخاب عمل متداولی مانند روش بولتزمن یا ε -greedy استفاده کرد. در حالت انتخاب عمل حریصانه، شرط زیر به هنگام انتخاب NA برقرار خواهد بود:

$$Q_i(s'_i, NA) > Q_i(s', a_i) \quad \forall a \in A_i. \quad (۴-۲۹)$$

برای به روز رسانی ارزش‌های حالت-عمل، می‌بایست تخمین‌ای از پاداش بازگشته انجام شود. اگر مساله تک-گام باشد، آن‌گاه پاداش بازگشته معادل با سیگنال تقویت لحظه‌ای خواهد بود. حالت دیگر وقتی است که پاداش بازگشته به صورت متوسط پاداش ایزود یا پاداش discount شده با افق نامحدود و یا از این دست باشد (به روابط ۴-۲، ۵-۲ و ۶-۲ مراجعه کنید).

در ابتدا حالت پاداش بازگشته‌ی معادل با پاداش لحظه‌ای را در نظر می‌گیریم. با توجه به تعریف $Q_i(s', a_i)$ و تعریف تحریک شدن و توجه به این نکته که پاداش بازگشته این حالت معادل با سیگنال تقویت لحظه‌ای $r(s)$ است، $Q_i(s', a_i)$ را می‌توان با در نظر گرفتن سیگنال تقویت لحظه‌ای که عامل دریافت می‌کند تخمین زد. اگر B_i رفتار کنترل‌کننده باشد و عمل a_i را انتخاب کرده باشد، یاری‌بهر آن رفتار به طور مستقیم در سیگنال تقویت دریافتی آن لحظه‌ی عامل متجلی می‌شود. پس با داشتن سیگنال تقویت آن لحظه می‌توان تخمین $Q_i(s', a_i)$ را به روز کرد. حال اگر B_j عمل مجازی‌ی NA را انتخاب کند، رفتار پایین‌تری چون B_i کنترل‌کننده می‌شود و سیگنال تقویت $r(s)$ را دریافت می‌کند. با استدلالی به مانند قبل، سیگنال تقویت نتیجه‌ی یاری‌بهر عمل انجام شده توسط B_i بوده است و در نتیجه یاری‌بهر رفتار کنترل‌کننده مشخص است. اما علاوه بر آن این سیگنال تقویت ناشی از سرکوب نکردن B_j و انتخاب NA نیز بوده است و میزان یاری‌بهر آن رفتار دقیقاً معادل با سیگنال تقویت آن لحظه خواهد بود. قواعد به روز رسانی در حالت پاداش لحظه‌ای بدین گونه است:

$$Q_{i_{k+1}}(s'_i, a_i) = (1 - \alpha_{k,i}(s'_i, a_i))Q_{i_k}(s'_i, a_i) + \alpha_{k,i}(s'_i, a_i)r(s) \quad (4-30)$$

(B_i is controlling behavior in s'_i and selects a_i)

$$Q_{j_{k+1}}(s'_j, NA) = (1 - \alpha_{k,j}(s'_j, NA))Q_{j_k}(s'_j, NA) + \alpha_{k,i}(s'_j, NA)r(s)$$

($\forall B_j; B_j \succ^T B_i$ and B_i is controlling behavior and B_j s are excited in s'_j and select NA)

$$(4-31)$$

توجه کنید که به روز رسانی NA می‌بایست برای همه‌ی رفتارهای تحریک‌شده‌ی بالاتر از رفتار کنترل‌کننده انجام شود. هم از تحلیل این بخش و هم از تحلیل بخش پیش بر می‌آید که رفتارهای پایین‌تر از رفتار کنترل‌کننده تغییری نمی‌کنند.

برای حالت پاداش بازگشته‌ی غیرلحظه‌ای (مثلا پاداش متوسط در طول اپیزود)، می‌توان از تخمین مونت‌کارلوی^{۱۹} پاداش بازگشته استفاده کرد. در این حالت از روش‌های bootstrapping مانند $TD(0)$ ([Sutton88]) به دلیل درست نبودن ارزش حالت بعدی به عنوان تخمین‌ای برای ارزش حالت فعلی نمی‌توان استفاده کرد.

دلیل این موضوع و عدم امکان استفاده از روش‌های bootstrapping تفاوت مناطق تحریک رفتارها و درک متفاوت آن‌ها نسبت به دنیا است. دو رفتار B_1 و B_2 را در نظر بگیرید. آن‌ها به کمک نگاشت‌های M_1 و M_2 دنیا را از طریق بازنمایی درونی‌ی حالت‌های‌شان مشاهده می‌کنند. تعریف $S^\circ(s'_i)$ را به یاد آورید (۴-۲۴). فرض کنید که دو حالت داخلی‌ی آن رفتارها به طور هم‌زمان در زیرمجموعه‌ی $S^\circ_1(s'_i) \cap S^\circ_2(s'_j)$ تحریک شده‌اند. برای راحت‌تر بودن تصور می‌توانید شکل (۴-۱) را ببینید و فرض کنید که درباره‌ی s'_5 از رفتار B_1 و s'_2 از رفتار B_2 صحبت می‌کنم. دیده می‌شود که آن‌ها با هم اندکی هم‌پوشانی دارند ولی قسمت‌های دیگری از آن‌ها جدا از هم است و در نتیجه وقتی یکی از آن‌ها تحریک می‌شود، دیگری ممکن است تحریک نشده باقی بماند. در حالت کلی $S^\circ_1(s'_i) \cap S^\circ_2(s'_j) \neq S^\circ_k(s'_i)$ ($k=1,2$). اینک تصور کنید که $s \in S^\circ_1(s'_i) \cap S^\circ_2(s'_j)$ و $B_1 \succ B_2$ و هم‌چنین B_1 - با توجه به $Q_1(s', \cdot)$ تصمیم می‌گیرد که NA انتخاب کند. در نتیجه B_2 که تحریک شده است عمل‌ای پیش‌نهاد می‌کند و کنترل عامل را در اختیار می‌گیرد و در این میان عامل پاداش‌ای (یا تنبیه) دریافت می‌کند. اگر قصد استفاده از bootstrapping را داشته باشیم می‌بایست از $Q_2(s'_2, a_i)$ برای کمک در تخمین $Q_1(s'_1, NA)$ (به عنوان تخمین فعلی از پاداش بازگشته شروع شده از حالت فعلی) استفاده کنیم. اما با این کار به خطا می‌رویم چون $Q_2(s'_2, a_i)$ متوسط پاداش بازگشته در زیرفضای مشترک $S^\circ_1(s'_i) \cap S^\circ_2(s'_j)$ نیست بلکه متوسط پاداش بازگشته در $S^\circ_2(s'_j)$ است. بدین ترتیب با انجام این کار تخمین نادرستی به دست آمده است. به همین دلیل برای تخمین پاداش

¹⁹ Monte Carlo

بازگشته هر حالت-عمل هر رفتار می‌بایست همه‌ی پاداش‌های دریافت شده از آن حالت-عمل را تا انتهای اپیزود نگهداری کرد و سپس با استفاده از روش مونت کارلو تخمین لازم را به دست آورد – مثلاً با استفاده از روش مونت کارلو اولین-ملاقات^{۲۰}.

اگر شخصی‌ای قصد پیاده‌سازی روش مونت کارلو را داشته باشد، احتمالاً بازده سیستم در مقایسه با مسایلی که در آن‌ها bootstrapping ممکن است، کندتر به حد مطلوب‌اش هم‌گرا می‌شود. با وجود این‌که این مساله ناشی از دشواری ذاتی یادگیری رفتار در معماری رفتارگرایی سلسله‌مراتبی‌ای است که هر کدام از رفتارهای‌اش نگاه‌ای مختص به خود به دنیای اطراف دارند، اما به نظر نمی‌رسد مشکل چندان محدودکننده‌ای باشد. بر این اعتقاد که در طراحی سیستم‌های رفتارگرایی فرض داشتن سیگنال تقویت لحظه‌ای کاملاً قابل قبول است. یک معماری رفتارگرایی معمولاً به عنوان کنترل‌کننده‌ی راکتیو عامل در نظر گرفته می‌شود و در این شرایط انتظار می‌رود که خوبی یا بدی و واکنش عامل بدون تاخیر مشخص شود و در نتیجه امید دریافت سیگنال تقویت لحظه‌ای برای حالت-عمل‌های مختلف درست همان چیزی است که می‌بایست بهینه شود. با این حال، چارچوب معرفی شده برای یادگیری رفتار سلسله‌مراتبی قابلیت کار با پاداش‌های بازگشته‌ی کلی را نیز دارد.

۴-۵) یادگیری هم‌زمان رفتار و ساختار

در دو فصل اخیر فرمول‌بندی مناسبی برای یادگیری ساختار و یادگیری رفتار معرفی کردم. ارزش کلی عامل (V_T) به اجزای ساده‌تری تقسیم شد و علاوه بر آن روش تخصیص اعتبار مناسبی برای یادگیری ساختار و یادگیری رفتار پیش‌نهاد کردم. طراح با توجه به مساله‌ی مورد نظر و دانش اولیه‌ای که در اختیار دارد می‌تواند از یادگیری ساختار، یادگیری رفتار و یا هر دوی آن‌ها به صورت هم‌زمان استفاده کند. استفاده‌ی هم‌زمان بدان معنا است که عامل با توجه به سیگنال تقویت دریافتی، نگاشت مناسب بین ادراک و عمل هر رفتار را بیابد و در همان هنگام ترتیب مناسب آن رفتارها در معماری را نیز بیاموزد (حل مساله‌ی مطرح شده در (۱۳-۲)). البته باید توجه شود که ارزش هر عمل هر رفتار ($Q_i(s_i, a_i)$) به مکان آن رفتار در ساختار نیز بستگی دارد. در صورت تغییر مکان رفتار در ساختار، دانش پیشین‌اش به

²⁰ First Visit Monte Carlo Method

- Initialize Learning parameters (different $\{\alpha\}$ depending on the method)
- For a lifetime do
 - Update learning parameters to ensure convergence
 - Select an architecture T^* that maximizes (3-6) (Zero Order representation) or (3-18) (First Order representation) using an optimization method (If the architecture is fixed, skip this step).
 - Using the provided architecture T^* , let the agent interact with the environment for a while
 - Receive reinforcement signal from the critic (that can be external or internal)
 - Update the estimation of necessary values ($\{\tilde{V}_{ij}\}$ for Zero Order structure learning representation (3-16), $\{\tilde{V}_{i>j}\}$ and $\{\tilde{V}_{i0}\}$ for First Order structure learning representation (3-28) and (3-29), and $\{Q_i(s_i, a_i)\}$ for Behavior learning (4-30) and (4-31) (immediate reward case) or Monte Carlo method for general return).

جدول ۴-۱. الگوریتم کلی برای نمایش عمل کرد عامل‌ای که ساختار و رفتارش را یاد می‌گیرد.

صورت دانش اولیه‌ای برای یادگیری شرایط جدید در می‌آید. در بسیار از تغییر مکان‌ها هم‌چنان اطلاعات مفیدی در ارزش حالت-عمل رفتارها باقی می‌ماند. با این وجود همیشه به هر رفتار اجازه‌ی تطبیق با شرایط جدید داده می‌شود و در نتیجه حتی اگر تخمین ارزش‌اش صحیح نباشد، رفتار می‌تواند با یادگیری مجدد آن را تصحیح کند. در این فرآیند تطبیق، لزومی به یک‌سان بودن سرعت یادگیری رفتارها و ساختار وجود ندارد؛ ساختار می‌تواند کندتر از رفتارها تغییر کند (با کم کردن نرخ یادگیری یا به روزرسانی‌ی ساختار در هر چند اپیزود یک‌بار). ایده‌ی پس این کار سریع‌تر کردن فرآیند تطبیق رفتارها در هنگام تغییر ساختار است. در نهایت، الگوریتم‌ای که نشان می‌دهد عامل یادگیر چگونه می‌بایست در طول دوره‌ی زندگی با کنش با محیطش ساختار و رفتارش را تغییر دهد در جدول ۴-۱ آمده است.

(۵) تکامل رفتار

۵-۱) مقدمه

در این فصل با رویکردی تکاملی و جمعیتی به مساله‌ی طراحی سیستم‌های رفتارگرا می‌پردازیم. برای شروع، اشاره مجدد به این نکته مفید است که طراحی یک سیستم را می‌توان به عنوان یک فرآیند جستجو در میان همه‌ی طراحی‌های ممکن در نظر گرفت. هدف طراحی نیز پیدا کردن عضوی از مجموعه‌ی فضای مساله^۱ است که معیاری از پیش تعیین شده را بهینه کند. در صورت پیدا کردن عضو مطلوبی از آن فضا، می‌توان ادعا کرد که به پاسخ مناسبی برای مساله‌ی طراحی رسیده‌ایم. روش‌های بسیاری به طور صریح یا ضمنی به بهینه‌سازی معیاری می‌پردازند. به عنوان مثال در شبکه‌های عصبی با سرپرست^۲ معیاری از خطای پیش‌بینی حداقل می‌شود. از همین دیدگاه بهینه‌سازی می‌توان به یادگیری تقویتی نگریست. یادگیری تقویتی معیاری از پاداش‌ها و تنبیه‌های دریافتی را حداکثر می‌کند. این معیار می‌تواند متوسط، متوسط وزن‌دار و یا مقدار لحظه‌ای پاداش‌ها و تنبیه‌های عامل باشد. این خاصیت یادگیری تقویتی را می‌توان به این صورت تفسیر کرد که یادگیری تقویتی تخمین‌ای از برنامه‌ریزی پویا^۳ است. با این دید، هدف فصل‌های پیش پیدا کردن پاسخ بهینه‌ای برای مساله‌ی طراحی عامل‌ای با معماری رفتارگرایی سلسله‌مراتبی با استفاده از فرمول‌بندی خاص یادگیری تقویتی بود.

برای هدف ما —طراحی عامل‌ای رفتارگرا— یادگیری تقویتی تنها و یا به‌ترین گزینه نیست. هم‌گرایی روش‌های متداول در یادگیری تقویتی در همه‌ی مسایل تضمین شده نیست. برای مسایل‌ای که از دید عامل مارکوف نباشد و یا از تقریب‌زن‌های عمومی‌ای^۴ چون شبکه‌های عصبی برای نگه‌داری تابع ارزش (یا ارزش-عمل) استفاده شده باشد، رسیدن به نقطه‌ی بهینه و یا حتی هم‌گرایی به پاسخ‌ای تضمین شده نیست. به عنوان نمونه‌ای از قیدهای لازم برای هم‌گرایی، می‌توان به Q-Learning اشاره

^۱ Problem Space

^۲ Supervised

^۳ Dynamic Programming

^۴ General Function Approximator

کرد که تنها در صورتی به پاسخ بهینه هم‌گرا می‌شود که مساله از دید عامل مارکوف باشد، بازنمایی تابع ارزش-عمل به صورت جدول ارجاعی و البته کاوش عامل نیز به میزان کافی صورت گرفته باشد (مثلا شرط $GLIE^5$ برقرار باشد [Jaakkola94]). مثال‌های ساده‌ای از مسایل غیرمارکوف وجود دارد که باعث واگرایی پروسه‌ی یادگیری می‌شوند. هم‌چنین مشکل‌دار بودن استفاده از تقریب‌زن‌های عمومی‌ای مانند شبکه‌های عصبی نیز کاملا شناخته شده است ([Thrun93] و [Boyan95]). یا به عنوان نمونه‌ای دیگر می‌توان به روش‌های سیاست‌محور^۶ی اشاره کرد که به طور مستقیم سعی در پیدا کردن بهینه‌ی سیاست عامل با استفاده از روش‌های مبتنی بر گرادیان می‌کنند (به عنوان مثال به [Peshkin01] مراجعه کنید). این روش‌ها که اخیرا بسیار در حل مسایل POMDP به کار می‌رود و نتایج موفق‌ای نیز داشته‌اند (مثلا [Kohl04]) با وجود آن‌که در طیف وسیع‌تری از شرایط هم‌گرا می‌شوند ولی تنها هم‌گرایی‌ی آن‌ها به بهینه‌ی محلی تضمین شده است. به همین دلایل است که توجه به روش‌های دیگری که بتوانند این مسایل را حل کند توجیه‌پذیر خواهد بود.

روش‌های تکاملی به عنوان روش جستجوی‌ای که از طبیعت الهام گرفته شده است، گزینه‌ی مطلوبی برای حل بعضی از این مشکلات است. استفاده از روش‌های تکاملی نه تنها از لحاظ زیست‌شناختی توجیه‌پذیر است، بلکه به دلیل قابلیت‌های ویژه‌ی آن‌ها گزینه‌ی بسیار مطلوبی به نظر می‌آید. به عنوان مثال، ویژگی‌های زیر از خواص مطلوب روش‌های تکاملی‌اند:

- جستجوی آن‌ها کلی^۷ است و نه محلی^۸؛ در نتیجه احتمال گیرافتادن در نقاط بهینه‌ی محلی کمتر است.
- نیازی به خوش‌تعریف بودن معیار مورد نظر ندارد؛ لزومی به مشتق‌پذیر بودن و یا حتی پیوسته‌بودن معیار مورد نظر نیست.

⁵ Greedy In Limit with Infinite Exploration

⁶ Policy-based

⁷ Global

⁸ Local

در نتیجه در صورت به کارگیری روش‌ای تکاملی برای جستجوی فضای مساله‌مان دیگر چندان نگران گیرافتادن در ماکزیمم‌های محلی تابع ارزش نخواهیم بود. همچنین به دلیل آن که در اکثر مسایل متداول برای عامل‌های در محیط واقعی قرارگرفته، تابع ارزش نسبت به پارامترهای توصیف‌کننده‌ی آن مشتق‌پذیر نیست، استفاده از روش‌ای که نیاز به وجود مشتق نداشته باشد بسیار مطلوب می‌نماید. البته در به کارگیری این روش‌ها نیز می‌بایست توجه کافی به خرج داد. روش‌های تکاملی نیز دارای مشکلاتی چون هم‌گرایی زود هنگام^۹ و یا گمراه‌کننده^{۱۰} بودن مساله دارند که باعث دشواری و یا نرسیدن به پاسخ مطلوب می‌شود. درست به همین دلیل است که پژوهش‌های بسیار زیادی در همین زمینه برای گسترش و تقویت روش‌های اولیه‌ی الهام گرفته شده از طبیعت در بین پژوهش‌گران آن رشته انجام می‌شود.

همان‌طور که در فصل اول توضیح دادم، استفاده از روش‌های تکاملی در روباتیک چندان تازه نیست. اکثر این روش‌ها سعی در پیدا کردن کنترلری یک‌پارچه برای تامین هدف مورد نظر دارند. این کنترل در اکثر موارد بوسیله‌ی شبکه‌ای عصبی توصیف می‌شود و نقش تکامل نیز پیدا کردن پارامترهای مطلوب است. تابع هدف می‌تواند از فرم‌های ساده‌ای مانند امتیاز منفی دادن به برخورد با اجسام تا فرم‌های پیچیده‌تری مانند امتیاز دادن به رفتار حرکت به سمت جسم‌ای مشخص باشد. با این‌همه یکی از بزرگ‌ترین ایرادهای روش‌های تکاملی فعلی، مقیاس‌پذیر^{۱۱} نبودن آن‌ها است. به کارگیری روش‌های تکاملی در مسایل اندکی پیچیده بسیار دشوار می‌شود. همچنین به دلیل تمایل کم و بیش همیشگی به بهره‌گیری از شبکه‌های عصبی به عنوان کنترلر روبات، قابلیت استفاده‌ی مجدد^{۱۲} از کنترلرهای تکامل یافته به صورت ماجول‌های از پیش موجود بسیار کم و یا حتی غیرممکن است. البته در این میان تلاش‌هایی برای حل این قبیل مشکلات و حل مسایل پیچیده‌تر با استفاده از روش‌های تکاملی انجام گرفته است که می‌توان به روی‌کرد ماجولار و روی‌کرد افزایشی^{۱۳} -افزایش تدریجی پیچیدگی تابع سازگاری- اشاره کرد [Togelius04].

⁹ Premature Convergence

¹⁰ Deceptive

¹¹ Scalable

¹² Reusability

¹³ Incremental

مطلوب است که متدولوژیی ما برای طراحی یک معماری رفتارگرا ویژگی‌های زیر را

تامین کند:

- به سرعت به پارامترهای کنترلر مطلوب هم‌گرا شود.
- سرعت واکنش‌اش به تغییرات سریع محیط بالا باشد.
- قابلیت کار در معماری‌ای سلسله‌مراتبی و ایجاد تصمیم‌گیری‌های چندسطحه را داشته باشد.
- طراحی به ماجول‌های مجزا و قابل استفاده‌ی مجددی بینجامد.

مکانیزم‌های متداول در روباتیک تکاملی عموماً می‌کنند. در اکثر آن‌ها کنترلر در طول دوره‌ی زندگی یک عامل ثابت است. پس در صورت تغییر محیط، کنترلر نمی‌تواند خود را با آن هم‌گام کند. ایده‌های یادگیری‌ای که در فصل‌های پیش مطرح شد چنین عیب‌ای ندارد ولی متأسفانه به دلایل توضیح داده شده - همیشه قادر به پیدا کردن بهترین کنترلر نیست. در نتیجه، استفاده‌ی هم‌زمان از یادگیری و تکامل - به طور هم‌زمان - جالب توجه می‌نماید. اضافه بر این، اگر ساختار تکاملی به گونه‌ای باشد که بتواند کنترلرهای سلسله‌مراتبی تولید کند، آن‌گاه با چارچوب مطرح شده در این پایان‌نامه سازگار خواهد بود و ویژگی‌هایی چون اولویت‌بندی، تصمیم‌گیری در سطوح مختلف تجرید زمانی و مکانی و ... را دارا خواهد بود. هم‌چنین در صورت ماجولار بودن، خاصیت استفاده‌ی مجدد نیز پیدا خواهد کرد.

با در نظر گرفتن این نکات، سعی در استفاده از ایده‌های تکاملی کردم. در فصل‌های پیشین مساله‌ی یادگیری را به دو بخش مجزای یادگیری ساختار و یادگیری رفتار تقسیم کردم. در یادگیری ساختار فرض بر این بود که تعدادی رفتار مناسب از پیش موجود است و هدف پیدا کردن ترتیب بهینه برای آن‌ها است. در یادگیری رفتار، ساختار از پیش فرض می‌شد و هدف به پیدا کردن نگاشت بین ورودی و خروجی هر رفتار فروکاسته می‌شد.

در این پژوهش تصمیم بر آن شد که تغییر رفتارها به عهده‌ی فرآیندی تکاملی باشد. به عبارت دیگر، به جای استفاده از یادگیری ساختار/یادگیری رفتار در معماری‌ای که همه‌ی اجزای آن تغییر

می‌کند از یادگیری ساختار/تکامل رفتار استفاده شود. علت این انتخاب این فرض است که یک عامل با داشتن مجموعه‌ای کافی از رفتارها و استفاده از سیگنال تقویت می‌تواند به‌ترین ساختار ممکن را به سرعت و در طی زندگی‌اش به دست بیاورد در حالی که پیدا کردن نگاشت مناسب هر رفتار به دلیل آن که فضای رفتارها معمولاً بزرگ‌تر از فضای ساختارها است، معمولاً کندتر است. فرض کنید ساختار T دارای m لایه باشد که می‌بایست یکی از n رفتار موجود در آن انتخاب شود. در قسمت‌های پیش داشتیم که کاردینالیت‌های فضای بازنمایی برای بازنمایی‌های یادگیری ساختارمان بدین‌گونه‌اند (تکرار ۳-۷ و ۳-۲۳):

$$\text{cardinality}(\text{ZO}) = n \cdot m \quad (۵-۱)$$

$$\text{cardinality}(\text{FO}) = n^2 + n \quad (۵-۲)$$

در مسایل متداول، n و m چندان بزرگ نیستند؛ مثلاً در مسالهی بلندکردن جسم به صورت چندروپاته که در فصل مربوط به آزمایش‌ها (فصل ۶) توضیح خواهیم داد هر دوی این‌ها ۵۸ و در نتیجه فضای جستجو خیلی بزرگ نیست. به همین دلیل می‌توان انتظار داشت که یادگیری ساختار بتواند در طول زندگی‌ی عامل به ساختار مناسبی دست‌پیدا کند. از طرف دیگر، در بسیاری از مسایل ممکن است ابعاد فضای جستجو برای رفتارها بسیار بزرگ باشد. به عنوان مثال رفتاری که در $|S|$ حالت مختلف واکنش نشان بدهد و دارای $|A|$ عمل مختلف است، دارای فضای جستجویی با بعد $|S| \cdot |A|$ خواهد بود. این که در مسایل عملی این مقادیر چقدرند موضوعی است که نیاز به بررسی‌ی چندین نوع مسالهی مختلف دارد اما من فرض را بر این گذاشتم که این بعد در اکثر مسایل بزرگ‌تر است و در نتیجه به احتمال زیاد مساله دشوارتر است. البته لازم به ذکر است که بزرگ‌بودن فضای بازنمایی مساله به تنهایی بیان‌گر میزان دشواری‌ی آن نیست. برای قضاوت در مورد پیچیدگی، ابتدا می‌بایست معیار پیچیدگی مشخص شود (مثلاً پیچیدگی‌ی تعداد نمونه‌های لازم برای یادگیری با دقت مشخص (پیچیدگی

نمونه^{۱۴}) اهمیت دارد یا مثلاً کیفیت به‌ترین پاسخ‌ای که می‌توان به آن رسید) که بررسی‌ی دقیقی‌ای در این‌باره صورت نگرفته است.

با این انتخاب و در صورت درست بودن فرض‌های انجام شده، به مکانیزمی می‌رسیم که ساختار معماری توسط یادگیری تغییر می‌یابد و مجموعه‌ی رفتارها نیز توسط تکامل. اما برای سازگار بودن با چارچوب مطرح در این پژوهش می‌بایست به این نکته نیز توجه کرد که لازم است تکامل توسط مکانیزم یادگیری‌ی ساختار قابل به کارگیری باشد. استفاده از ایده‌های متداول در اکثر پیاده‌سازی‌های رباتیک تکاملی – که به صورت کنترلر شبکه‌ی عصبی‌ای یک‌پارچه پیاده‌سازی می‌شوند – قابل استفاده نیست. عناصر اولیه یادگیری‌ی ساختار، مجموعه‌ای از رفتارها $\{B_i\}$ هستند و نه کنترلری که ترکیب همه‌ی این‌ها با هم باشد. به همین دلیل، تصمیم گرفتیم تا روش مورد نظر به جای تکامل کل رفتار عامل با هم، به تکامل تک‌تک رفتارها به طور جداگانه بپردازد. در نتیجه حاصل تکامل نه یک کنترلر یک‌پارچه که مجموعه‌ای از رفتارهای مجزا خواهد بود. به همین دلیل ایده‌های سلسله‌مراتبی مطرح شده به همان شکل قبل قابل به کارگیری خواهد بود و همچنین طراح در نهایت به مجموعه‌ای از رفتارها می‌رسد که که قابلیت استفاده‌ی مجدد دارند. بدین ترتیب، ایده‌ی کلی‌ی استفاده از تکامل برای ایجاد رفتارهای مناسب ایجاد شد. در بخش بعد، به توضیح جزئیات این ایده می‌پردازیم.

۵-۲) تکامل رفتارها

فرض کنید که می‌خواهیم معماری رفتارگرایی برای عامل بسازیم. برای این کار می‌بایست مجموعه‌ی مناسب‌ای از رفتارها $\{B_i\}$ و همچنین ساختار مناسب T را بیابیم به طوری که اهدافمان برآورد شود. اگر هدف را به صورت تابع تقویت $R: S \times a \times S' \mapsto \mathbb{R}$ (یا هر فرم مشابه دیگر) در نظر بگیریم، آن‌گاه هدف‌مان بیشینه کردن این تابع – یا تابع‌ای از این تابع^{۱۵} – خواهد بود. تعریف V_T (۷-۲) را به یاد بیاورید. هدف طراحی‌ی یک معماری، پیدا کردن ساختار و رفتارهای مناسب است (۱۳-۲). مشاهده

¹⁴ Sample Complexity

¹⁵ در شرایطی که از حالت‌های کلی‌ی پاداش بازگشته استفاده می‌کنیم که خود تابع‌ای از سیگنال تقویت است.

می‌شود که دو مجموعه پارامتر مختلف برای بهینه‌سازی وجود دارد: پارامترهای بیان‌گر هر رفتار و پارامتر بیان‌گر ساختار.

اگر روش‌های حل مساله را به دو دسته‌ی روش‌های یادگیر و روش‌های تکاملی تقسیم کنیم، آن‌گاه یکی از رویکردهای زیر را می‌توان برگزید:

- یادگیری ساختار – یادگیری رفتار
- تکامل ساختار – یادگیری رفتار
- تکامل ساختار – تکامل رفتار
- یادگیری ساختار – تکامل رفتار
- یادگیری و تکامل ساختار – یادگیری و تکامل رفتار (و ترکیب‌های مشابه)

البته از ذکر حالت‌های بدیهی‌ای مانند ثابت فرض کردن ساختار و یادگیری (یا تکامل) رفتار و برعکس بدلیل آن‌که تنها سعی در حل بخشی از کل مساله می‌کنند صرف‌نظر شده است. همان‌طور که در مقدمه گفته شد، به نظر می‌رسد در اکثر موارد رفتارها آن بخشی‌ای از مساله‌اند که تغییرات آن‌ها باعث تغییرات کم‌تری در رفتار کلی‌ی عامل می‌شود در حالی که تغییر کوچک‌ای در ساختار (جابه‌جا شدن دو رفتار یا اضافه شدن یک رفتار جدید) باعث تغییرات بسیار در رفتار کلی‌ی عامل می‌شود. هم‌چنین به دلیل آن‌که در اکثر مسایل فضای جستجوی رفتارها بسیار بزرگ‌تر از فضای جستجوی ساختار است، ایجاد تغییرات مورد نظر در سطح رفتارها نیاز به تغییرات بسیار بیش‌تری نسبت به تغییر در ساختار دارد و در نتیجه سرعت واکنش عامل به تغییرات محیط بسیار کم‌تر خواهد بود. به عبارت دیگر ساختار بخشی‌ای از عامل است که تغییرات سریع از آن انتظار می‌رود و رفتارها بخش کندتر آن هستند. به همین دلیل، تصمیم بر آن شد که رفتارها تکامل یابند و ساختار یادگرفته شود. توضیح یادگیری ساختار در فصل‌های پیش آمده است و دوباره تکرار نمی‌شود. اینک به تکامل رفتار می‌پردازم.

۵-۲-۱) تکامل هم‌کارانه‌ی رفتارها^{۱۶}

برخلاف اکثر روش‌های متداول تکامل در رباتیک، در روش‌ای که به کار گرفته‌ام کل کنترلر یک‌جا تکامل نمی‌یابد بلکه کنترلر عامل به اجزای رفتاری تقسیم شده و هر رفتار به طور جداگانه تکامل می‌یابد. رفتارها با توجه به نقش از پیش مشخص‌شده‌شان جداگانه در مخزن ژنتیکی^{۱۷} خویش تکامل می‌یابند و تمام رفتارهای مرتبط –مانند انتخاب به‌ترین‌های هر جمعیت و تزویج و جهش ژنتیکی– در همان بستر و جدا از دیگر رفتارها انجام می‌شود. به عنوان مثال اگر در مساله‌ی بلند کردن جسم توسط چندروبات، دو رفتار بلند کردن و تنظیم زاویه‌ی جسم وجود داشته باشد دو جمعیت مختلف از رفتارها خواهیم داشت که رقابت بین اعضای آن و عملیات ژنتیکی‌شان جداگانه انجام می‌شود. طراح از پیش مشخص می‌کند که بازنمایی هر رفتار به چه صورت است، سپس جمعیت‌ای از اعضایی با همان بازنمایی برای هر رفتار ایجاد می‌کند و پس از آن هر جمعیت با توجه به سازگاری‌ی اعضای‌اش عملیات ژنتیکی را انجام می‌دهد.

البته هنوز نکات بسیاری می‌بایست مطرح شود تا این مکانیزم قابل پیاده‌سازی باشد:

- چگونه با استفاده از این جمعیت‌های رفتاری می‌توان یک عامل خاص را به وجود آورد؟
- ارتباط سازگاری‌ی عامل و سازگاری‌ی هر کدام از رفتارهای تشکیل‌دهنده‌اش چگونه است؟
- اپراتورهای ژنتیکی هر رفتار چگونه هستند؟

نحوه‌ی ارتباط این سیستم تکامل هم‌کارانه‌ی چندجمعیتی و عامل که الهام گرفته شده از Enforced Sub-Population (ESP) است [Gomez97] بدین‌صورت است: پس از ایجاد تعدادی جمعیت برای هر رفتار، یکی از رفتارهای هر کدام از جمعیت‌ها را به طور اتفاقی انتخاب می‌کنیم. در نتیجه اگر n جمعیت B_i داشته باشیم، مجموعه‌ای n تایی $\{B_i\}$ خواهیم داشت که در آن $B_i \in B_j$. این مجموعه به عامل داده می‌شود و عامل با استفاده از این مجموعه می‌تواند با محیط‌ش ارتباط برقرار

¹⁶ Cooperative Evolution of Behaviors

¹⁷ Genetic Pool

کند. اگر عامل از یادگیری ساختار بهره ببرد می‌تواند به‌ترین ترتیب این رفتارها را بیابد و اگر نه از ساختار از پیش تعیین‌شده‌ای استفاده می‌کند. بدیهی است که کیفیت کار عامل به کیفیت هر کدام از رفتارها بستگی دارد: اگر رفتارهای "مناسب"ی در اختیار عامل باشد، عامل می‌تواند بازده بالایی داشته باشد (یا به عبارت دیگر، رفتار کلی‌ی سازگارانه‌تری نسبت به محیط داشته باشد)؛ ولی اگر رفتارهای "مناسب"ی در اختیارش قرار داده نشود، عامل حتی با وجود یادگیری ساختار نیز نمی‌تواند به بازده بالایی برسد. می‌توان نتیجه گرفت که کیفیت رفتار کلی‌ی عامل به کیفیت هر کدام از رفتارهایش ارتباط دارد. البته این که میزان یاری‌بهر^{۱۸} هر رفتار در رفتار کلی‌ی عامل چگونه است مشخص نیست – بعضی از رفتارها ممکن است تاثیر بیش‌تری داشته باشند و بعضی تاثیر کم‌تری، یا حتی بعضی رفتارها ممکن است بتوانند اثر سوء رفتارهای دیگر را بپوشانند یا برعکس. در این پژوهش دو مکانیزم مختلف برای ارزیابی‌ی تاثیر هر رفتار پیش‌نهاد شده است که اینک آن‌ها را توضیح خواهیم داد.

۵-۲-۲) اشتراک سازگاری یک‌نواخت

از آن‌جا که کیفیت و سازگاری‌ی عامل به سازگاری هر کدام از رفتارهای تشکیل‌دهنده‌اش بستگی دارد، می‌توان سازگاری‌ی هر کدام از آن رفتارها را متوسطی از سازگاری‌ی همه‌ی عامل‌هایی دانست که آن رفتار در آن‌ها نقش داشته است. به عبارت دیگر اگر نمونه‌ای^{۱۹} خاص از یک جمعیت مانند B_i^j (عضو j ام از جمعیت i ام) در p عامل مختلف قرار بگیرد و آن عامل‌ها دارای سازگاری‌های $f_i (i: 1 \dots p)$ باشند، آن‌گاه سازگاری‌ی این نمونه‌ی خاص رفتار به صورت زیر خواهد بود:

$$f^u(B_i^j) = \frac{1}{p} \sum_{i=1}^p f_p \quad (۵-۳)$$

برای f_i می‌بایست به کیفیت کار عامل در دوره‌ی زندگی‌اش توجه کرد. این کیفیت کار به میزان پاداش و تنبیه‌ای که دریافت می‌کند بستگی دارد. اگر طراح بخواهد همان معیاری را که در طول یادگیری سعی

¹⁸ Contribution

¹⁹ Instance

در بیشینه کردن‌اش داریم با روند تکامل بهینه کند، آن‌گاه می‌بایست سازگاری‌ی عامل را نیز از همان پاداش‌ها و تنبیه‌ها استخراج کند.^{۲۰} برای این منظور می‌توان نوشت:

$$V_{\{B\}}|_{\text{Last K episodes}} = f_{\{B\}}|_{\text{Last K episodes}} = E \left[\frac{1}{K} \sum_{t \in \text{Last K episode}} r_t \mid \text{the agent with } \{B\}, t \in \text{Last K episode} \right] \quad (۵-۴)$$

که بیان‌گر میزان دریافتی سیگنال تقویت در K اپیزود آخر است. مقدار K هر چقدر کوچک‌تر باشد نشان‌دهنده‌ی تمایل بیش‌تر به کیفیت نهایی‌ی عامل است و هر چقدر بزرگ‌تر باشد -تا حدی که در نهایت کل اپیزودهای دوره‌ی زندگی‌ی عامل را در بر بگیرد- بیش‌تر به کیفیت عامل در کل دوره‌ی زندگی‌اش نزدیک می‌شود. در این صورت سازگاری‌ی هر نمونه‌ی رفتار در اشتراک سازگاری‌ی یک‌نواخت به صورت زیر تعریف می‌شود:

$$f^u(B_i^j) = \frac{1}{N} \sum_{\{B\}_i} V_{\{B\}_i} |_{\text{Last K episodes}} \quad (۵-۵)$$

که $\{B\}_i$ ها N مجموعه رفتار تصادفی انتخاب شده‌ای هستند که $B_i^j \in \{B\}_i$. از آن‌جا که $\{B\}_i$ ها به طور تصادفی انتخاب شده‌اند ممکن است بعضی از رفتارهای جمعیت بیش از بقیه شانس شرکت در مجموعه‌ی رفتارهای عامل را داشته باشند. اما در صورتی که تعداد آزمون^{۲۱} ها زیاد باشد می‌توان مطمئن بود که هر رفتار به میزان کافی در عامل‌هایی با ترکیب رفتارهای مختلف شرکت جسته و تخمین به دست آمده معیار قابل قبولی از سازگاری‌ی آن رفتار است.

²⁰ البته طراح می‌تواند معیارهای دیگری را نیز در نظر بگیرد. آن‌گاه مساله‌ی بهینه‌سازی‌ای که بخش‌ای از آن (ساختار) در فرآیند یادگیری حل می‌شود دیگر با مساله‌ی بهینه‌سازی‌ای که بخش دیگر آن (رفتارها) توسط تکامل حل می‌شود یک‌سان نخواهد بود. با این‌حال طراح می‌تواند معیارهای دیگری مانند یادگیری‌ی سریع‌تر و ... را نیز در چارچوب معرفی شده در این فصل در نظر بگیرد.

²¹ Trial

۵-۲-۳) اشتراک سازگاری ارزش محور

اشتراک سازگاری یک نواخت بر پایه‌ی این فرض بود که سازگاری‌ی عامل و سازگاری‌ی هر کدام از رفتارهای درونی‌اش هم‌بسته‌اند و می‌توان سازگاری‌ی هر رفتار را میانگین‌ای از سازگاری‌ی همه‌ی عامل‌هایی دانست که آن رفتار در آن‌ها دخیل بوده است. می‌توان گفت رگه‌هایی از دیدگاه روان‌شناسی رفتارگرایانه در آن نوع اشتراک سازگاری وجود دارد که تلاش‌ای بر کنکاش درونی‌ی ذهن عامل ندارد. اما اگر به نحوی بتوان به دانش درونی‌ی عامل پی برد و میزان یاری‌بهر هر رفتار در سازگاری‌ی عامل را تشخیص داد، آن‌گاه ممکن است به نتایج به‌تر و دقیق‌تری در اشتراک سازگاری برسیم. خوش‌بختانه این اطلاعات در بازنمایی‌ای که استفاده می‌کنیم وجود دارد. بازنمایی مرتبه صفر را در نظر بگیرید. یاری‌بهر هر رفتار B_i برای عامل معادل کل ارزش‌ای است که در لایه‌های مختلف ساختار کسب کرده است. یعنی:

$$V_{T(B_i)} = \sum_{q=1}^m P\{B_i | L_q\} V_{qi} P(L_q \text{ is controlling}) = \sum_{q=1}^m \tilde{V}_{qi} \quad (5-6)$$

پس می‌توان یاری‌بهر هر رفتار در معماری‌ی T را از روی داده‌های بازنمایی مرتبه صفر تخمین زد. با توجه به این‌که هر نمونه‌ی رفتار B_i^j در چندین و چند معماری به کار می‌رود، عاقلانه است که سازگاری‌ی آن را به مانند روش پیشین با متوسط‌گیری بر روی چند رفتار محاسبه کنیم. به عبارت دیگر با همان تعاریف و علایم می‌توان نوشت:

$$f^v(B_i^j) = \frac{1}{N} \sum_{\{B_i\}} V_{T(B_i)} \quad (5-7)$$

توجه کنید که تعاریف به کاررفته سخت و انعطاف‌ناپذیر نیستند. مثلاً می‌توان K در (۵-۴) را به گونه‌ای عوض کرد که همه‌ی اپیزودها را شامل شود و معیار سازگاری، معیار سازگاری‌ی کل دوره‌ی زندگی عامل شود و یا این‌که آن را به چند اپیزود آخر محدود کرد تا معیاری از عمل‌کرد نهایی عامل باشد.

مشابه همین انعطاف برای اشتراک سازگاری ارزش محور نیز برقرار است. به عنوان نمونه اگر مقدار $\alpha_{n,ij}$ در به روزرسانی $\tilde{V}_{ij}$ به صورت $\alpha_{n,ij} = 1/n$ باشد، آن گاه $\tilde{V}_{ij}$ متوسط ارزش رفتار در لایه L_j خاص از معماری خواهد بود و متوسط گیری مطابق (۵-۶) برابر متوسط یاری بهر رفتار در طول زندگی عامل خواهد بود. هم چنین می توان تخمین ارزش های مرتبه صفری که برای یادگیری ساختار به کار می رود متفاوت با تخمین ای که برای محاسبه سازگاری هر رفتار به کار می رود انتخاب کرد.

۵-۲-۴) اپراتورهای ژنتیکی

انتخاب اپراتورهای ژنتیکی وابسته به نوع بازنمایی دانش ای اند که برای حل مساله به کار می رود. برای شبیه سازی های این بخش از بازنمایی جدول ارجاعی استفاده کرده ام. بازنمایی داخلی رفتار B_i (و در نتیجه B_i^j که نمونه ای از آن خانواده رفتار است) بدین گونه است:

$$B_i(s') : S'_i \rightarrow A'_i. \quad (۵-۸)$$

در نتیجه با دانستن این نگاشت حالت به عمل و دانستن حالت فعلی سیستم، عمل مناسب (یا NA) قابل انتخاب است. به مانند قبل هر رفتار متناسب با حالت سیستم می تواند تحریک شده یا نشده و متناسب با نوع نگاشت اش (انتخاب NA و یا غیر از آن) فعال شود یا نشود. این که کدام رفتار کنترل عامل را در اختیار بگیرد بستگی به فعال شدن یا نشدن بقیه رفتارهای معماری دارد. بُعد این نگاشت متناسب با بُعد حالت و بُعد خروجی رفتار است که توسط طراح مشخص می شود. در پیاده سازی هایی که انجام داده ام هر کدام از B_i^j به صورت ماتریس چندبعدی باینری ای ذخیره شده اند که فعال بودن بیت ای در آن معادل انتخاب عمل و فعال نبودن آن معادل انتخاب NA در آن حالت است. اپراتورهای ژنتیکی متناسب با این ساختار ویژه انتخاب شده اند. در این پژوهش دو نوع اپراتور ژنتیکی جهش و ترویج مخصوص فرم ماتریسی بازنمایی رفتار معرفی شده است.

اپراتور جهش به دو گونه‌اند: سخت و نرم. جهش سخت که با احتمال $p_{m_{hard}}$ انتخاب می‌شود، یک ماتریس B_i^j را با ماتریس تصادفی کاملاً جدیدی جای‌گزین می‌کند. جهش نرم اختلال‌ای^{۲۲} در هر ماتریس ایجاد می‌کند که این اختلال معادل جستجویی در اطراف پاسخ پیشین در فضای پاسخ‌ها است. در پیاده‌سازی باینری انجام گرفته، این اختلال با عوض کردن چند بیت در بازنمایی رفتار صورت می‌گیرد.

اپراتور تزویج بدین‌صورت تعریف شده است:

$$B_i^{j_{new}} = XB_i^{j_{old}} + \bar{X}B_i^{k_{old}} \quad (۵-۹)$$

که در آن X ماتریس تصادفی‌ای با ابعاد مناسب است که یک بودن درایه‌های آن حاصل آزمون تصادفی مستقل از هم با توزیع یک‌سان^{۲۳} برنولی^{۲۴} با احتمال p_c است؛ $P(X_{ij...l} = 1) = p_c$.

۵-۳) الگوریتم ممیتیک^{۲۵}

در بخش پیش چارچوب‌ای کلی مبتنی بر یادگیری و تکامل برای ایجاد خودکار معماری‌ای رفتارگرا پیش‌نهاد کردم. تکامل به عنوان جزء کند ولی همه‌جانبه‌نگر به جستجوی فضای بزرگ رفتارها می‌پردازد و یادگیری به عنوان جزء سریع ولی محدودتر ترتیب مناسب رفتارها را در ساختار می‌یابد. می‌توان این مجموعه ساز و کار را تقلیدی از طبیعت دانست که در آن هم تکامل و هم یادگیری در تطبیق موجودات با محیطشان نقش دارند.

در مکانیزم مطرح شده پس از پایان عمر موجود، همه‌ی تجربیات‌اش از بین می‌رود و تنها تاثیر یادگیری بر سازگاری‌اش باقی می‌ماند. هر موجود جدیدی که در محیط شروع به فعالیت می‌کند بدون هیچ تجربه‌ی قبلی‌ای و تنها با اتکا با ذخیره‌ی ژنی‌اش شروع به یادگیری و کنش با محیط می‌کند.

²² Perturbation

²³ Independent Identical Distribution (i.i.d.)

²⁴ Bernoulli

²⁵ Memetic Algorithm

استفاده‌ی بیش‌تر از تجربیات موجود در طول زندگی و انتقال آن به نسل‌های بعد وسوسه‌برانگیز می‌نماید. هدف این بخش پیش‌نهاد روش‌ای تازه برای استفاده از این دانش است. به دلیل شباهت کلی‌ی روش با ایده‌های مطرح در زمینه‌ی پژوهشی‌ی الگوریتم‌های ممیتیک، نام این بخش را چنان انتخاب کردم، اگرچه وجود نام مشترک الزاما به معنای تفسیر دقیقا یکسان نیست.

الگوریتم‌های ممیتیک [Moscato92] بر پایه‌ی ایده‌ی مم^{۲۶} بنا شده‌اند. مم واحد اطلاعات‌ای است که به هنگام ارتباط مردم با هم بازتولید می‌شود [Dawkins76]. اما با وجود تعریف گسترده‌ی مم، در جامعه‌ی محققان محاسبات تکاملی، الگوریتم ممیتیک اغلب به صورت ترکیب‌ای از روش‌ای تکاملی و جستجویی محلی دانسته می‌شود. بدین نحو که با استفاده از روش‌های مختلف جستجوی محلی که برای مساله‌ی مورد نظر طراحی شده است نتیجه‌ی حاصل از فرآیند تکامل را بهبود می‌بخشند. این نوع پیاده‌سازی‌ی مم حاصل برداشت‌ای از مم است که در آن بخش تکاملی‌مانند الگوریتم، دیگر فرآیندی ژنتیکی نیست بلکه شبیه‌ساز بستر جابه‌جایی مم‌ها است و جستجوی محلی نیز معادل با فرآیند یادگیری و تطبیق‌ای است که هر فرد جامعه بر واحدهای اطلاعاتی‌ای که با آن سر و کار دارد اعمال می‌کند. پس روش‌های متداول در تکامل در این‌جا بر فرهنگ‌ها و ایده‌ها اعمال می‌شوند و نه بر ژن‌ها. بهبود فردی‌ای که به عنوان بخش‌ای از الگوریتم ممیتیک صورت می‌گیرد و معمولا به صورت جستجوی محل‌ای پیاده‌سازی می‌شود ممکن است در مراحل مختلفی از فرآیند بهینه‌سازی رخ دهد: پیش از محاسبه‌ی سازگاری و بدون تاثیر بر آن‌چه مشابه با ژنوتایپ^{۲۷} در الگوریتم‌های تکاملی است، پس/پیش از اعمال اپراتورهای تکاملی و ... با این دیدگاه، الگوریتم‌های ممیتیک نام دیگر ترکیب‌های جستجوی محلی/تکاملی‌ای مانند تکامل بالدوینی^{۲۸} و تکامل لامارکی^{۲۹} است (و البته اکثر پیاده‌سازی‌های الگوریتم‌های ممیتیک به تکامل لامارکی – تکامل‌ای که در آن ژنوتایپ نیز تغییر می‌کند – شبیه است). با این‌که از دید نویسنده، این نوع برداشت – یعنی ترکیب جستجوی محلی در بین عملیات ژنتیکی – از ایده‌ی مم و الگوریتم ممیتیک محدودکننده است اما همین نوع استفاده نیز به نتایج قابل توجه‌ای دست

²⁶ Meme

²⁷ Genotype

²⁸ Baldwinian Evolution

²⁹ Lamarckian Evolution

یافته است و هم‌چنان به عنوان یکی از راه‌های پر فایده و آینده‌دار محاسبات تکاملی شناخته می‌شود. برای مشاهده‌ی نمونه‌هایی از پژوهش‌های انجام شده درباره‌ی الگوریتم‌های ممیتیک به [Moscato92]، [Radcliffe94]، [Merz99]، [Merz00]، [Federici03]، [Buriol04]، [Krasnogor04]، [Ong04] و [Zou04] مراجعه کنید. برای آشنایی بیش‌تر با الگوریتم‌های ممیتیک و پژوهش‌های انجام شده می‌توانید به [Moscato03] نگاه کنید.

در پژوهش حاضر، دو تعبیر مختلف از الگوریتم ممیتیک شده است. تعبیر اول همان استفاده‌ی توام تکامل و یادگیری برای حل مسالهی طراحی رفتار و ساختار است که جزییات آن تاکنون معرفی شده. تفاوت اصلی‌ی این تعبیر با آنچه در بدنه‌ی اصلی‌ی الگوریتم‌های ممیتیک استفاده می‌شود در اختلاف سوژه‌ی بهینه‌سازی در دو مرحله‌ی یادگیری و تکامل است. در اکثر الگوریتم‌های ممیتیکی، بهینه‌سازی مرحله‌ی تکامل و مرحله‌ی جستجوی محلی (یادگیری) بر روی یک فضا صورت می‌گیرد در حالی که در این پژوهش جستجوی تکامل بر فضای رفتارها است و جستجوی یادگیری بر فضای ساختارها. این نوع به کارگیری جدا از تفاوت یاد شده مشابه با روش تکاملی‌ی بالدوینی است.

اما تعبیر استعاره‌ی مم به صورت‌های دیگری نیز ممکن است. به عنوان مثال مم می‌تواند به صورت سنت فرهنگی برای حل یک مسالهی مشترک برای اعضای جامعه نیز تعبیر شود. این سنت می‌تواند در شکل‌دهی‌ی باور افرادی که تازه به آن فرهنگ پا گذاشته‌اند —مثلاً تازه به دنیا آمده‌اند— نقش داشته و آن‌ها را به سمت آنچه پیش از آن تجربه شده و آزمون پس داده ره‌نما باشد. در این تعبیر مم چیزی جدا از ژن موجودات است و در سطح‌ای دیگر عمل می‌کند. مم به صورت جهت‌ده^{۳۰} اولیه و به عنوان نقطه‌ی آغازین‌ای برای جستجوی پاسخ مناسب مساله عمل می‌کند. برای روشن شدن قضیه این دو ایده‌ی به کارگیری مم را در قالب مسالهی اصلی‌ی این پایان‌نامه مطرح می‌کنم.

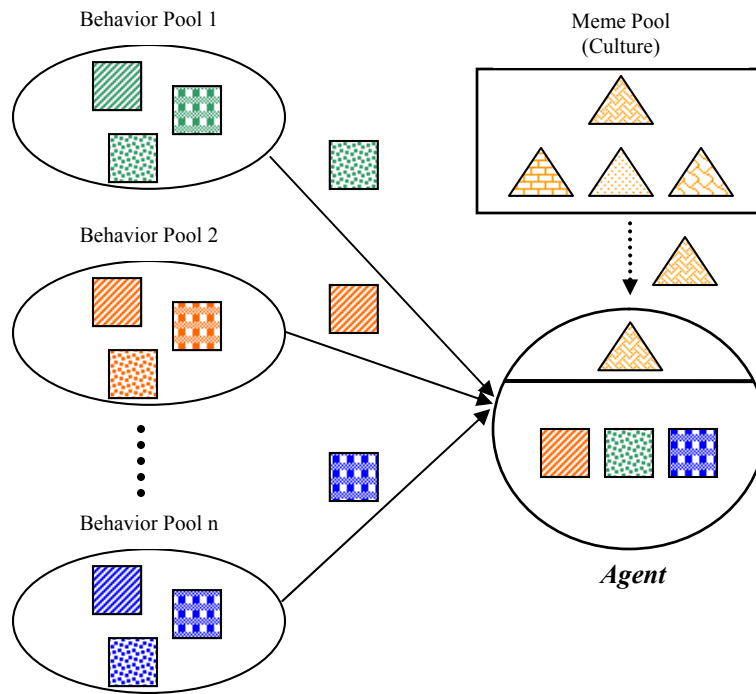
مکانیزم کلی‌ی طراحی سیستم رفتارگرای مان بدین‌گونه خواهد بود: چندین نوع رفتار مختلف هر کدام در مخزن ژنتیکی‌ی خویش به طور جداگانه تکامل می‌یابند. در این سیستم تکامل هم‌کارانه، هر مخزن ژنتیکی دارای ویژگی‌ها و ژنوتایپ خاص خود است. طبیعت به هنگام ایجاد موجودی (عامل) جدید یکی از اعضای هر کدام از این مخزن‌ها را به طور اتفاقی انتخاب می‌کند تا در نهایت مجموعه‌ای

³⁰ Bias

از رفتارهای $\{B_i\}$ ایجاد شود. عامل تازه به دنیا آمده با استفاده از این رفتارها (اطلاعات ژنتیکی) با کنش با محیط اطرافاش و دریافت سیگنال تقویت سعی در بهینه‌سازی معیار عمل کرد خویش - و در نتیجه یادگیری به‌ترین ساختار قابل ساخت با آن مجموعه رفتارها- می‌کند. پس از آن میزان سازگاری عامل به طبیعت برگردانده می‌شود و طبیعت با توجه به روش اشتراک سازگاری، سازگاری هر کدام از رفتارها را مشخص می‌کند. در این روند، ترکیب تکامل و یادگیری تا حدی مشابه با تفسیر غالب از الگوریتم‌های ممیک است. اما برخلاف اکثر روش‌های الگوریتم ممیک‌ای، یادگیری (و بهینه‌سازی) به طور مستقیم بر اطلاعات ژنتیکی (یعنی $\{B_i\}$) اعمال نمی‌شود بلکه بر روی ساختاری انجام می‌شود که اطلاعات ژنتیکی به عنوان اجزای آن (رفتارها) به کار می‌رود. به بیان‌ای دیگر فضای مساله‌ی یادگیری و تکامل - با وجود ارتباط تنگاتنگ با هم - یک‌سان نیستند. نحوه‌ی پیاده‌سازی برداشت دوم از مم - یعنی برداشت مم به عنوان سنت و باور اولیه - بدین‌صورت است: عامل در انتهای زندگی‌اش به‌ترین ساختار یافته‌شده را به عنوان مم به فرهنگ‌اش^{۳۱} (مخزن مم^{۳۲}) باز می‌گرداند. با توجه به میزان سازگاری عامل، شانس باقی ماندن و غلبه پیدا کردن مم در آن فرهنگ کم یا زیاد می‌شود. ساختارهایی که باعث ایجاد عامل‌های موفق شده‌اند، مم‌های غالب پیدا می‌کنند و ساختارهای غیرموثر و ناموفق قادر به ایجاد سنت باقی نخواهند شد. به هنگام تولید موجودی جدید، مخزن مم به عنوان باور اولیه‌ای (دانش اولیه‌ای) درباره‌ی ساختار مطلوب عمل می‌کند و عامل به جای شروع از ساختاری تصادفی، از آن ساختار از-پیش جواب‌داده استفاده می‌کند. اگر ساختار مناسب موجود فعلی دقیقاً همان‌ای باشد که موجودات پیشین به آن رسیده بودند (مثلاً به دلیل شباهت مجموعه رفتارهای در اختیارشان)، آن‌گاه عامل از همان اولین قدم‌های زندگی‌اش (اپیزودهای آغازین) درست و صحیح عمل می‌کند. البته از آن‌جا که رفتارها در طول نسل‌های متوالی تغییر می‌کنند، فرهنگ نسل‌های پیش ممکن است به‌ترین انتخاب برای نسل فعلی نباشد. متناسب با میزان تغییر اعضای پیشین جامعه نسبت به اعضای فعلی، مم‌های فرهنگ یا کاملاً مناسب یا نسبتاً مناسب خواهند بود. این نکته وقتی بیش‌تر آشکار می‌شود که به

³¹ Culture

³² Meme Pool



شکل ۵-۱. ساخت عامل جدید با استفاده از اطلاعات ژنتیکی و اطلاعات فرهنگی

خاطر بیاوریم تغییر رفتارها (تغییرات ژنتیکی) معمولاً آهسته و پیوسته است. نمایش‌ای از این فرآیند در شکل ۵-۱ نشان داده شده است.

تاکید دوباره بر این شکل که در فرآیند مطرح شده از دو تفسیر مختلف الگوریتم ممیک استفاده شده است مفید خواهد بود: تفسیری که الگوریتم ممیک معادل با جستجوی پس از عملیات تکامل است (استفاده از یادگیری) و تفسیری که در آن از مهم‌ها به عنوان دانش اولیه‌ی عامل‌های جدید استفاده شده است. نحوه‌ی پیاده‌سازی برداشت اول در بخش پیشین این فصل توضیح داده شده است. هدف اصلی‌ی این بخش معرفی مکانیزم‌ای برای تفسیر دوم است. شبه‌کدی از این فرآیند را در جدول ۵-۱ می‌بینید.

بیان دقیق‌تر ایده و چگونگی پیاده‌سازی الگوریتم ممیک در مساله‌ی خودکارسازی طراحی عامل رفتارگرا بدین گونه است: مجموعه‌ی عامل‌های $\{agent_i\}$ را در نظر بگیرید. هر عامل در این پژوهش شامل دو بخش‌اند: رفتارها و ساختار آن؛ $agent_i : (\{B_i\}, T_i)$. هر عامل در دوره‌ی زندگی‌اش تلاش می‌کند با تغییر ساختار، حداکثر پاداش ممکن را دریافت کند (به یاد داشته باشید که در

این بخش تنها ساختار یاد گرفته می‌شود و رفتارها ذخیره‌ی ژنتیکی به ارث رسیده هستند. گرچه اعمال ایده‌های این بخش به حالت‌ای که رفتار نیز یاد گرفته شود چندان متفاوت نیست). در انتهای زندگی عامل، نتیجه‌ی یادگیری $(\{B_i\}, T_i^*)$ خواهد بود که در آن T_i^* ساختار بهینه شده است. این ساختار بهینه اینک به صورت مم‌ای وارد فرهنگ (یا همان مخزن مم) $\mathbf{M}$ می‌شود. در این مخزن، مم‌های عامل‌های سازگارتر شانس غلبه‌گی‌ی بیش‌تری دارند. این مم‌ها پایدارترند و شانس بیش‌تری برای تاثیرگذاری بر عامل‌های تازه به دنیا آمده دارند. به عبارت دیگر، فرهنگ $\mathbf{M}$ به صورت مجموعه‌ای از ترکیب‌های زیر است:

$$\mathbf{M} : \{ (T_i^*, f_{T_i}) \} \quad (5-10)$$

که در آن f_{T_i} سازگاری‌ی عامل‌ای با ساختار (T_i) است. بدیهی است که ساختار T_i می‌تواند نتیجه‌ی نهایی‌ی یادگیری عامل‌هایی با مجموعه رفتارهای متفاوت باشد - یعنی ما به ازای هر T_i چندین مجموعه $\{B_i\}$ داریم. مشخص است که به ازای رفتارهای مختلف، سازگاری‌ی عامل نیز تغییر می‌کند. در نتیجه برای این که معیاری حدودی از میزان ارزش هر مم داشته باشیم می‌بایست متوسط ارزش آن را به ازای همه‌ی عامل‌هایی که ساختارشان در نهایت به آن مم هم‌گرا خواهد شد محاسبه کنیم، یعنی:

$$f_{T_i} = \frac{1}{N} \sum_{A_i: (\{B_i\}, T_i)} f(A_i) \quad (5-11)$$

که در آن N تعداد عامل‌هایی است که در نهایت و پس از یادگیری به معماری T_i رسیده‌اند. اضافه کردن فاکتور فراموشی^{۳۳} اجازه‌ی تحمل نایستایی^{۳۴} فرآیند تکاملی را می‌دهد:

³³ Forgetting Factor

³⁴ Non-Stationary

$$f_{T_{i+1}} = (1 - \alpha_{T_i})f_{T_{i_n}} + \alpha_{T_i}f(A) \quad A: (\{B_i\}, T_i) \quad (5-12)$$

اندازه‌ی مخزن مم محدود است: تنوع فرهنگی^{۳۵} یک اجتماع بی‌نهایت نیست. در نتیجه می‌بایست راه‌کار معین‌ای برای حفظ و نگه‌داری‌ی مم‌ها داشته باشیم. فرض کنید مم‌ای جدید که پیش از آن عضوی از مخزن نبوده است پس از دوره‌ی زندگی موجودی ایجاد شود. اگر مخزن مم‌ها گنجایش داشته باشد، اضافه کردن مم تازه به آن مشکل‌ای ایجاد نمی‌کند. اگر مم‌ای ایجاد شده قبلاً هم در مخزن وجود داشته بود، آن‌گاه به کمک معادله‌ی (۵-۱۲) ارزش آن مم در مخزن مم‌ها به روز می‌شود. اما اگر مم در مخزن وجود نداشت و از طرف دیگر مخزن نیز پر شده بود (ظرفیت مخزن را طراح از پیش مشخص می‌کند)، مم جدید جای‌گزین بدترین مم - مم با پایین‌ترین ارزش - می‌شود. چنین فرآیندی در جوامع بشری نیز رخ می‌دهد: الگوهای رفتاری جدید به تدریج جای‌گزین سنت‌های قدیمی و بی‌فایده می‌شوند. هم‌چنین در طول زمان ارزش یک سنت با توجه به سودی که می‌رساند کم یا زیاد می‌شود.

پس از تولد هر موجود، مم‌ای از مخزن مم‌ها به عنوان دانش اولیه به او منتقل می‌شود. این که کدام مم انتخاب شود بستگی به ارزش هر کدام از مم‌ها دارد: مم‌های با ارزش‌تر شانس بیش‌تری برای باور آغازین موجودات شدن دارند.

در این پژوهش از مکانیزم رقابتی‌ای استفاده کرده‌ام که شانس انتخاب هر مم متناسب با ارزش مقیاس‌شده‌ی^{۳۶} هر کدام از مم‌ها است. چنین فرآیندی در روش‌های تکاملی‌ی متداول نیز به کار می‌رود. هم‌چنین گاهی مم‌ای به طور اتفاقی - بدون توجه به ارزش‌اش - انتخاب می‌شود. این کار تا حدی مشابه با اپراتور جهش در روش‌های تکاملی، باعث حفظ تنوع فرهنگی جامعه می‌شود. پس از انتخاب مم با هر روش ممکن، عامل از ساختار مشخص شده توسط مم با محیط کنش می‌کند و با استفاده از روش یادگیری ساختاری که در اختیار دارد سعی در بهبود ساختار خویش و یافتن ساختار بهینه (البته مقید به

³⁵ Cultural Diversity

³⁶ Scaled

رفتارهای در اختیارش) می‌کند. اگر رفتارهای عامل تا حدی قابل قبولی شبیه به رفتارهایی باشند که عامل تولید کننده‌ی مم اولیه‌ی عامل اخیر از آن استفاده می‌کرده است، آن‌گاه عامل از همان ابتدا با ساختاری تقریباً درست و صحیح شروع به تصمیم‌گیری می‌کند. در این صورت دوره‌ی لازم برای یادگیری - و رسیدن به حد قابل قبول‌ای از بازده - بسیار کم خواهد شد.

البته اضافه کردن جهت‌دهی اولیه ممکن است باعث اضافه شدن احتمال گیر افتادن عامل در بیشینه‌گی محلی‌ای شود: تجربیات "خوب" - ولی نه "خیلی خوب" - موجودات قبلی باعث جلوگیری از کاوش و تجربه‌ی حالت‌های "خیلی خوب" می‌شود. البته کاوش‌گری در انتخاب ساختارها - و مثلاً تجربه‌ی گاه به گاه ساختارهای تصادفی - و همچنین ایجاد گاه‌گاه مم‌های اتفاقی باعث کم شدن احتمال این اتفاق می‌شود.

با توجه به آن‌چه تاکنون در این فصل گفته شد، در چارچوب مطرح شده سه روش مختلف برای تطبیق عامل‌ها در نظر گرفته شده است: تکامل ژنتیکی که در آن مجموعه‌ای از رفتارها در مخزن‌های ژنتیکی جداگانه‌ای تکامل می‌یابند، یادگیری فردی هر عامل برای یافتن بهترین ساختار ممکن و جهت‌دهی فرهنگی‌ای که از تجربیات عامل‌های پیشین برای راهنمایی یادگیری عامل‌های جدید استفاده می‌شود.

- Initialize n different behavior pools $\{B_i\}$ for each different behavior type B_i
- Initialize an empty culture (meme pool) $\mathbf{M}$
- While stopping condition are not met
 - Selects n different behavior B_i from each behavior pool to make a set of randomly chosen behaviors $\{B_i\}$
 - If there is any meme in the meme pool $\mathbf{M}$, select a meme $T^0 \in \mathbf{M}$ according to the merit of each meme
 - Pass $\{B_i\}$ and T^0 to the agent
 - Initialize Learning parameters
 - For a lifetime do
 - Update learning parameters to ensure convergence
 - Select an architecture T^* that maximizes (3-6) (Zero Order representation) using an optimization method (If the architecture is fixed, skip this step).
 - Using the provided architecture T^* , let the agent interact with the environment for a while
 - Receive reinforcement signal from the critic (that can be external or internal)
 - Update the estimation of necessary values ($\{\tilde{V}_{ij}\}$ for Zero Order structure learning representation (3-16))
 - Return Fitness and the final structure T^* to the evolutionary process
 - Share fitness to behaviors according to sharing policy ((5-5) or (5-7))
 - Update meme pool using T^* and its fitness (5-12)
 - For each behavior pool
 - Apply conventional genetic operators to behaviors in order to generate a new population, i.e. Selection, Crossover, and Mutation.

جدول ۵-۱. الگوریتم ممینک. رفتارها تکامل می‌یابند، ساختار یاد گرفته می‌شود و مم‌ها انتقال دانش ساختار مناسب را بر عهده دارند.

۶) آزمایش‌ها

در این فصل به بررسی تجربی روش‌های پیشنهاد شده برای یادگیری و تکامل معماری رفتارگرایی سلسله مراتبی می‌پردازیم. برای این کار آن روش‌ها را در دو مساله‌ی مختلف آزمایش می‌کنیم. یکی از مسایل، مساله‌ای مجرد است که به طور خاص برای تست روش‌های پیشنهاد شده طراحی شده است. طراحی چنین مساله‌ای ما را قادر می‌سازد تا در شرایطی کنترل‌شده به بررسی روش‌ها بپردازیم. مساله‌ی دیگر، مساله‌ای از دنیای واقعی است؛ بلندکردن یک جسم به وسیله‌ی چند روبات. پیش از پژوهش اخیر، این مساله‌ی دشوار به وسیله‌ی معماری‌ای رفتارگرا با سعی و خطای بسیار طراح حل شده بود [Nili01]. از آن‌جا که پروسه‌ی طراحی کنترلر مناسب برای این مساله ساده نبوده است، طراحی خودکار معماری رفتارگرایی که بتواند آن مساله را حل کند اثبات تجربی قابل قبولی برای موثر بودن روش‌های پیشنهادی است.

بررسی‌های انجام گرفته را به دو بخش روش‌های مبتنی بر یادگیری و روش‌های مبتنی بر ترکیب یادگیری/تکامل تقسیم می‌کنیم. در بخش یادگیری به یادگیری ساختار با رفتار از پیش مشخص، یادگیری رفتار با ساختار معلوم، و یادگیری هم‌زمان ساختار و رفتار می‌پردازیم. در بخش ترکیب یادگیری/تکامل به تکامل رفتار با ساختار از پیش مشخص و تکامل رفتار و یادگیری ساختار به طور هم‌زمان می‌پردازیم. همچنین ایده‌ی الگوریتم ممیتیک معرفی شده در فصل تکامل رفتار (فصل ۵) نیز هم‌زمان با ایده‌های تکامل رفتار مورد مطالعه قرار می‌گیرد.

پیش از ادامه لازم است بگوییم که هیچ تلاش منظم و دقیقی برای بهینه‌سازی پارامترها یادگیری یا تکامل صورت نداده‌ام و در نتیجه ممکن است با انتخاب مجموعه پارامترها یا سیگنال تقویت دیگری – همان‌طور که در جای لازم بیش‌تر در موردش توضیح خواهم داد – به نتایج به‌تری رسید.

۶-۱) مساله مجرد

هدف مساله‌ی مجردی نشان دادن قابلیت عامل برای تطابق با محیطی ناشناخته و اتفاقی طراحی شده از طریق کنش با آن و دریافت سیگنال پاداش یا تنبیه است. این تطبیق می‌تواند به صورت یادگیری یا

تکامل باشد. اگر عامل بتواند سیاست مناسب‌ای برای کنش با محیط پیدا کند پاداش می‌گیرد و در غیر این صورت تنبیه می‌شود. تنبیه و یا پاداش عامل به تطبیق دقیق و مو به موی سیاست عامل و سیاست صحیح و دل‌خواه $a^* = \Pi(s)$ بستگی دارد. در این مساله $s \in \{1, 2, \dots, r\} \times \{1, 2, \dots, s\}$ و $a^* \in \{1, 2, \dots, q\}$ و $\Pi(\cdot)$ نیز در ماتریس‌ای است که نگاشت بین حالت فعلی و عمل مطلوب ولی ناشناخته برای عامل را مشخص می‌کند. عامل می‌بایست از طریق کنش با محیط سیاست $a = \pi^{agent}(s)$ را بیابد که کم‌ترین فاصله با a^* را داشته باشد. فاصله در فضای عمل‌ها بدین‌گونه تعریف می‌شود:

$$\|\Pi(s) - \pi(s)\| = \begin{cases} 0 & \text{if } \Pi(s) = \pi(s) \\ 1 & \text{otherwise} \end{cases} \quad (6-1)$$

و کل فاصله برای همه‌ی حالت‌ها نیز بدین‌گونه مشخص می‌شود:

$$\|\Pi - \pi\| = \sum_{s \in S} \|\Pi(s) - \pi(s)\|. \quad (6-2)$$

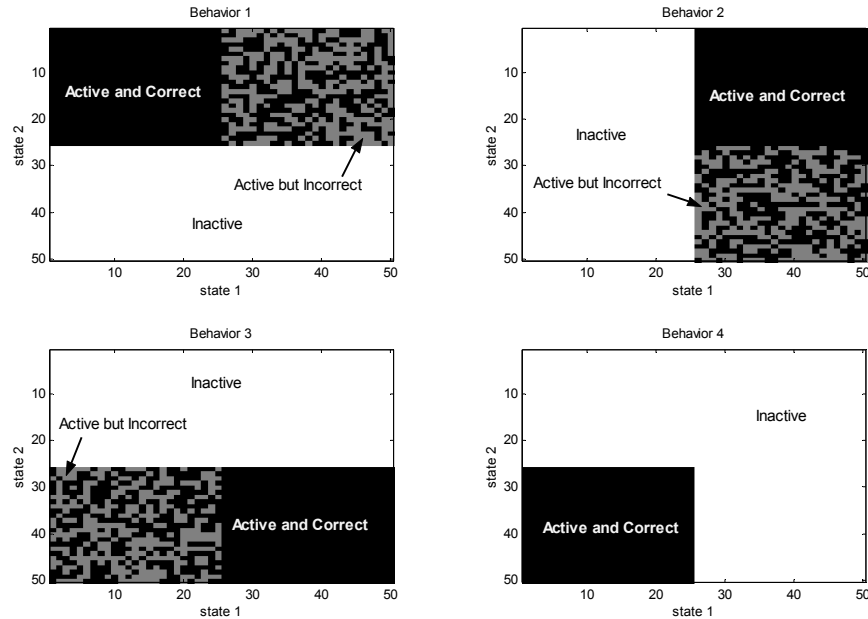
اندازه‌ی فاصله‌ی بین سیاست یادگرفته شده و سیاست بهینه (احتمال خطا) بدین‌گونه تعریف می‌شود:

$$P(error) = \frac{1}{|S|} \|\Pi - \pi\| \quad (6-3)$$

که همان‌طور که می‌بینید اثر ابعاد فضای حالت در آن با نرمال کردن خنثی شده است. این اندازه، معیار خوبی برای بازده سیستم در حالت‌ای است که توزیع انتخاب حالت‌ها یک‌نواخت باشد.

سناریوی کنش عامل با محیط و یادگیری بدین‌گونه است: عامل با معماری تصادفی T - که در آن رفتارها از جعبه‌ابزاری از رفتارها انتخاب شده است - آغاز به کار می‌کند. عامل در حالت

¹ در این‌جا فرض کرده‌ام که عامل اعمال‌اش را حریصانه انتخاب می‌کند و π^{agent} نیز سیاست deterministic عامل است.



شکل ۶-۱. مساله‌ی مجرد) چهار رفتار مختلف مساله‌ی مجرد. به مناطق مختلف تحریک و درستی یا نادرستی پاسخ در آن مناطق دقت کنید.

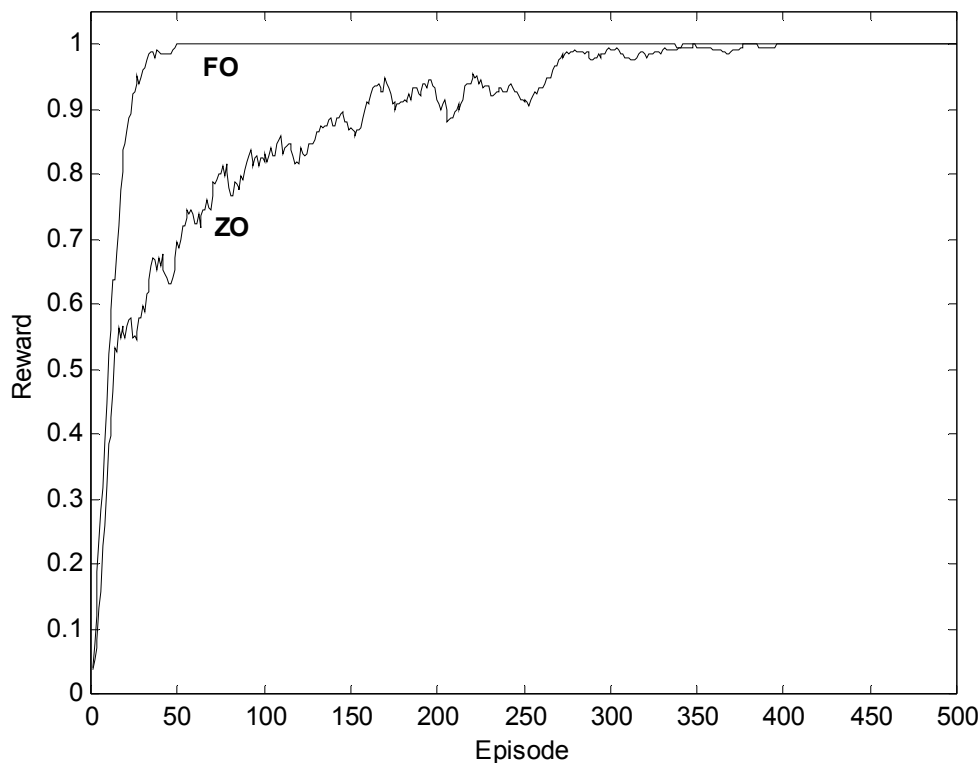
تصادفی s قرار می‌گیرد و سپس عمل a را با توجه به سلسله مراتب رفتارها و سیاست هر کدام از رفتارهای اش $B_i(s_i)$ انتخاب می‌کند. اگر عامل عمل صحیح در مقایسه با پاسخ بهینه $\Pi(s)$ را انتخاب کند، پاداش $+1$ دریافت می‌کند و در غیر این صورت به اندازه 1 - تنبیه می‌شود؛ یعنی:

$$r(s) = \begin{cases} +1 & \text{if } \Pi(s) = \pi^{agent}(s) \\ -1 & \text{otherwise} \end{cases} \quad (6-4)$$

که در آن عمل انتخاب شده‌ی عامل با توجه به سیاست هر رفتار و ساختار رفتارها در معماری است. این عمل با انتخاب s کاملاً تصادفی جدیدی ادامه می‌یابد. در مساله‌ی مجرد هر کنش با محیط و دریافت سیگنال تقویت یک اپیزود را تشکیل می‌دهد.

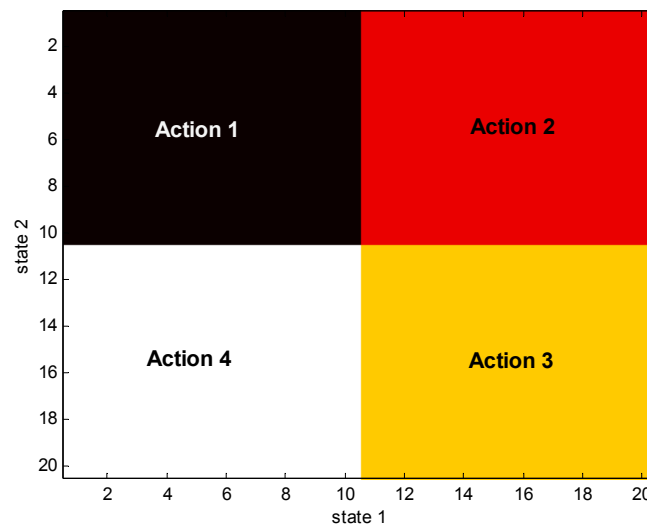
۶-۱-۱) یادگیری ساختار

اولین آزمون روش‌های یادگیری ساختار بر روی مساله‌ی مجرد انجام می‌شود. برای این کار مساله‌ای مجرد با ابعاد فضای حالت 50×50 و دارای دو عمل تعریف شده است (برای مشاهده‌ی جزئیات آزمون



شکل ۲-۶. (مساله‌ی مجرد) مقایسه بین روش‌های یادگیری ساختار (مبتنی بر بازنمایی مرتبه صفر و بازنمایی مرتبه یک)

به جدول ۱-۶ مراجعه کنید). چهار رفتار مختلف (B_1, B_2, B_3 و B_4) تعریف شده‌اند که هر کدام در بخش‌ای از فضای حالت تحریک می‌شوند (شکل ۱-۶). در شکل، ناحیه‌ی تیره بیان‌گر ناحیه‌ای از فضای حالت است که رفتار تحریک شده و عمل‌ای را -درست یا نادرست- پیش‌نهاد می‌کند. نواحی‌ای که عمل انتخاب شده درست است با رنگ سیاه و نواحی‌ای که عمل انتخاب شده نادرست است با رنگ خاکستری مشخص شده است. رفتارها به گونه‌ای طراحی شده‌اند که مساله‌ی چینش رفتارها در ساختار، دارای یک پاسخ بهینه‌ی عمومی یکتا باشد ($T = [B_1 \ B_2 \ B_3 \ B_4]$ در این مساله). در چنان حالت‌ای، معماری به ازای همه‌ی حالت‌ها به درستی پاسخ می‌دهد. نتیجه‌ی به کارگیری یادگیری ساختار مبتنی بر بازنمایی مرتبه صفر و مرتبه یک در شکل ۲-۶ نشان داده شده است. دیده می‌شود که هر دو روش به خوبی عمل کرده و عامل توانسته ساختار مناسب را به سرعت یاد بگیرد. با این وجود دیده می‌شود که روش یادگیری ساختار مبتنی بر بازنمایی مرتبه یک سریع‌تر هم‌گرا شده است. این پدیده به آن دلیل است که در روش مبتنی بر بازنمایی مرتبه یک، در هر کنش عامل با محیط، اطلاعات بیشتری

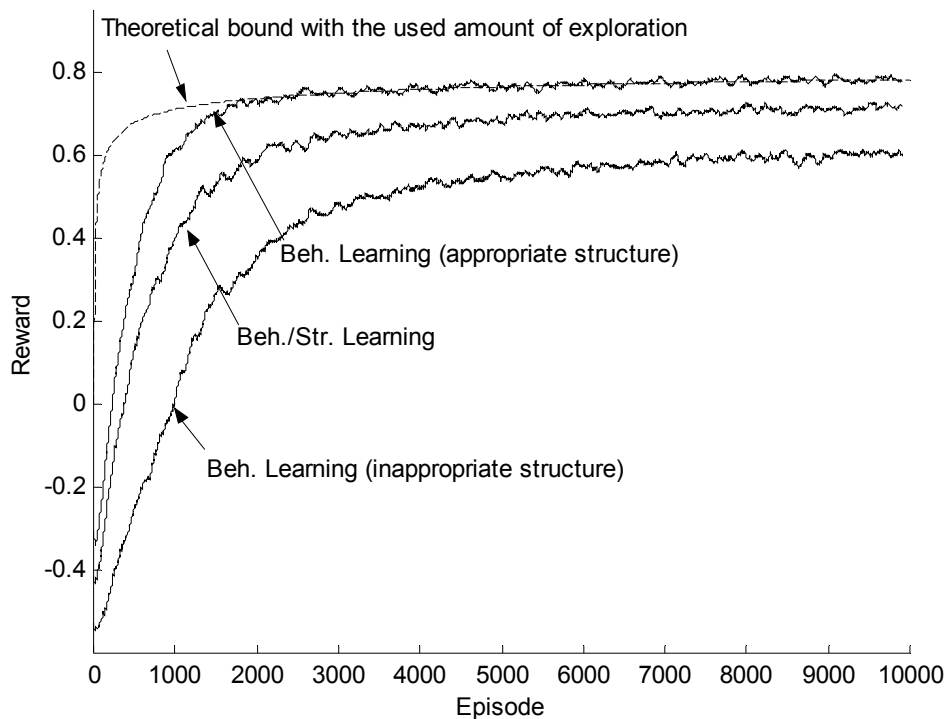


شکل ۳-۶. (مساله‌ی مجرد) فضای مساله برای آزمون‌های یادگیری رفتار/ساختار

درباره‌ی ساختار صحیح به شکل تقویت ترتیب رفتارها کسب می‌شود. همچنین از شکل مشخص است که روش بازنمایی مرتبه یک دارای نویز کم‌تری در میزان سیگنال تقویت دریافتی نسبت به بازنمایی مرتبه صفر است. این نویز و تاخیر همچنین به این خاطر است که در بعضی از تجربه‌های یادگیری با استفاده از بازنمایی مرتبه صفر، یادگیری دیرتر از معمول صورت می‌گیرد و آن تجربه‌های تاخیردار باعث کند شدن متوسط کلی‌ی بازنمایی مرتبه صفر می‌شود و گرنه در اکثر موارد یادگیری مرتبه صفر نیز به سرعت (سریع‌تر از آنچه در شکل دیده می‌شود) به پاسخ مناسب هم‌گرا می‌شود.

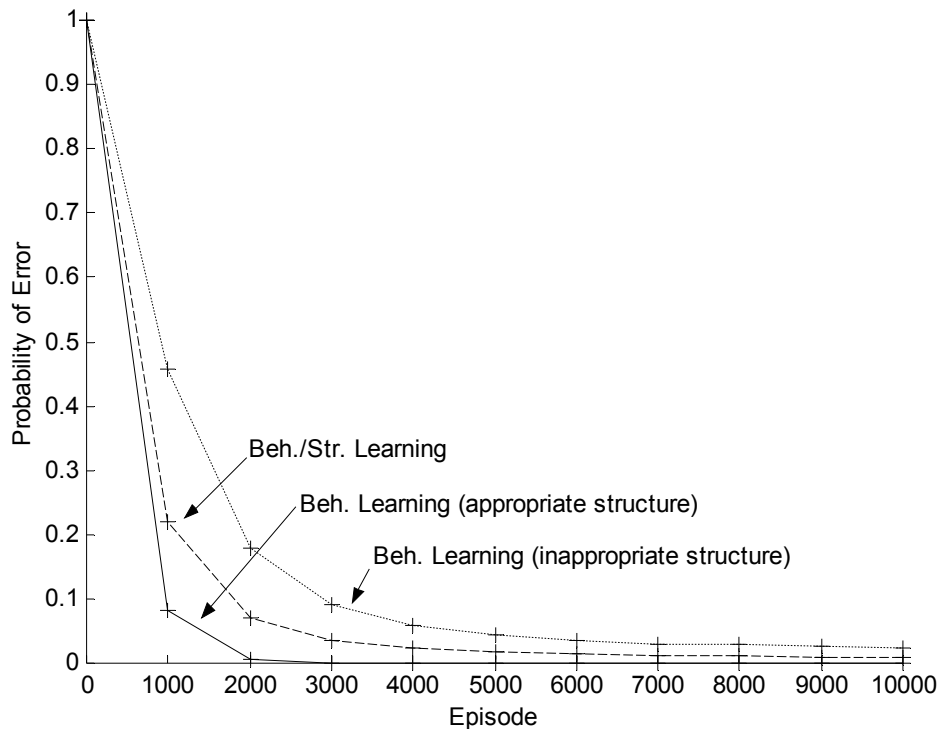
۶-۱-۲) یادگیری رفتار و یادگیری هم‌زمان رفتار/ساختار

در این زیربخش روش پیش‌نهادی برای یادگیری رفتار و همچنین یادگیری ساختار و یادگیری رفتار به طور هم‌زمان برای مساله‌ی مجرد مورد آزمون واقع می‌شود. فضای مساله $\Pi(\cdot)$ از چهار ناحیه تشکیل شده است که در هر کدام از آن نواحی فقط یکی از چهار عمل ممکن قرار گرفته است (شکل ۳-۶). مطابق قبل چهار رفتار داریم که نواحی‌ی تحریک‌شان به مانند آزمایش یادگیری ساختار تعریف شده است ولی برخلاف آن، رفتارها از پیش دارای عمل‌ای مشخص نیستند و عمل منسوب به هر حالت هر رفتار می‌بایست یاد گرفته شود. برای رفتار B_i فضای عمل توسعه یافته شامل عمل a_i و NA است. به دلیل ساختار ویژه‌ی مساله و نواحی‌ی تحریک رفتارها، همیشه نگاشت‌ای مناسب از حالت-عمل رفتارها



شکل ۴-۶. (مساله‌ی مجرد) مقایسه پاداش دریافتی برای یادگیری رفتار (ساختار ثابت) و یادگیری هم‌زمان ساختار و رفتار.

وجود دارد که منجر به تولید پاسخ صحیح معماری شود. با این وجود، با توجه به ساختار معماری، نداشت‌های متفاوتی می‌بایست یاد گرفته شوند. اگر ساختار مناسب باشد (به این معنا که هیچ رفتاری برای پاسخ صحیح معماری لازم به یادگیری NA نباشد) آن‌گاه یادگیری تنها یک نداشت از یک رفتار برای پاسخ مناسب دادن در هر حالت کافی است. اگر ساختار نامناسب باشد نه تنها نداشت رفتار پیش‌نهاد دهنده‌ی عمل صحیح می‌بایست به درستی یاد گرفته شود بلکه رفتارهای بالایی آن رفتار کنترل‌کننده نیز باید یاد بگیرند تا در آن حالت NA پیش‌نهاد کنند. برای این مساله، ساختار کاملاً مناسب $T = [B_1 \ B_2 \ B_3 \ B_4]$ و ساختار کاملاً نامناسب $T = [B_4 \ B_3 \ B_2 \ B_1]$ خواهد بود. انتظار داریم که ساختار مناسب با القای دانش اولیه‌ی صحیح، یادگیری را ساده‌تر و سریع‌تر و ساختار نامناسب آن را دشوارتر و کندتر کند.

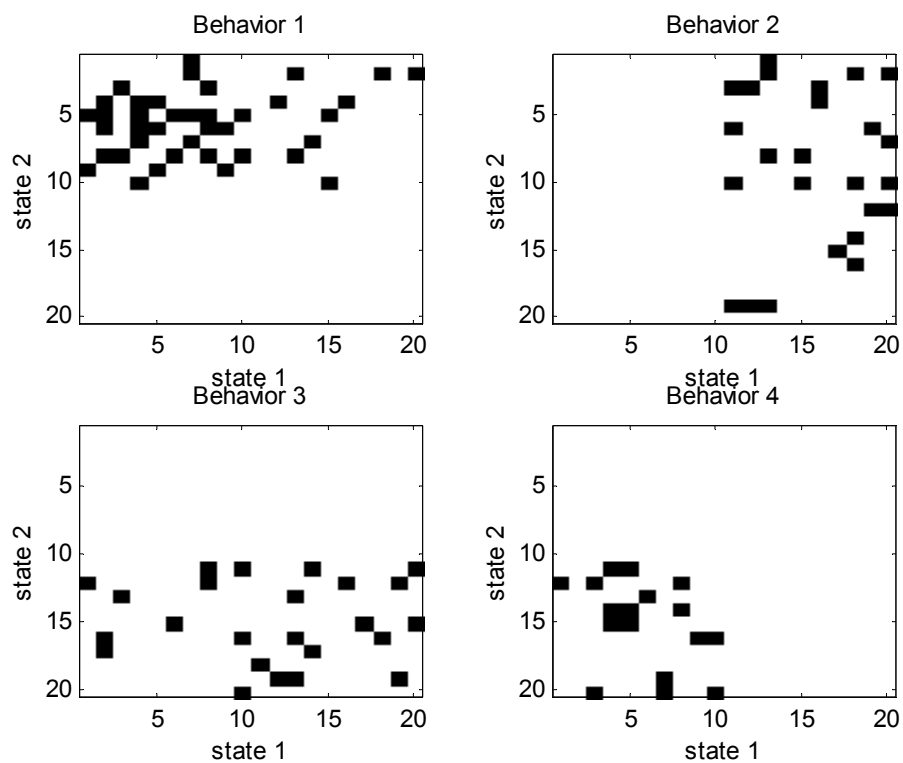


شکل ۵-۶. (مساله‌ی مجرد) مقایسه احتمال خطای یادگیری رفتار (ساختار ثابت) و یادگیری هم‌زمان ساختار و رفتار.

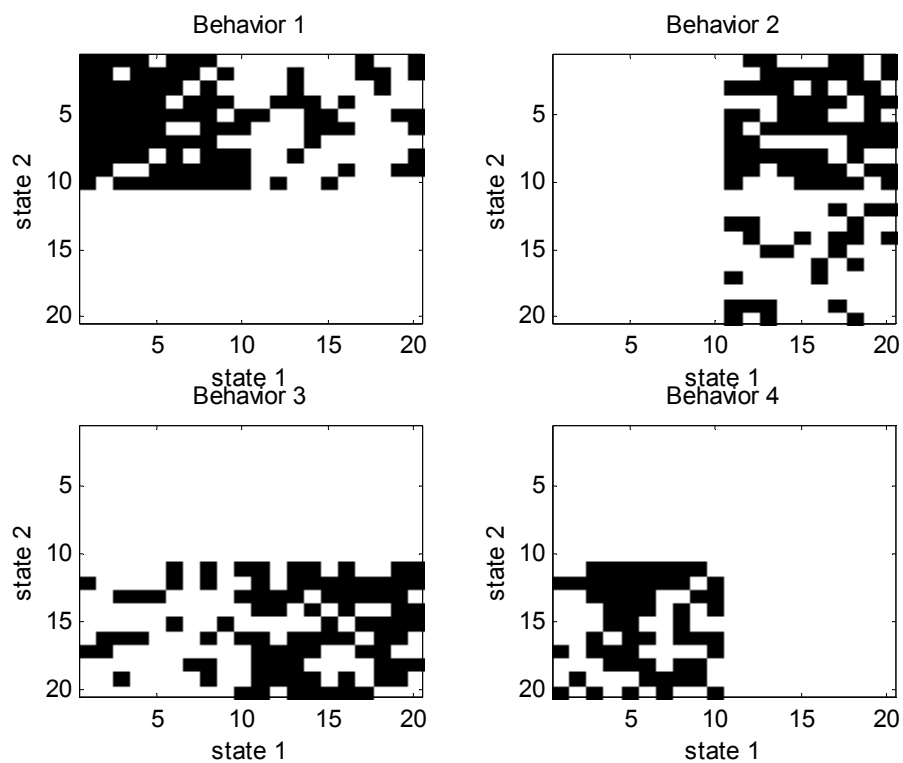
سه آزمون مختلف انجام می‌دهم: یادگیری رفتار با ساختار مناسب، یادگیری رفتار با ساختار نامناسب و یادگیری هم‌زمان ساختار و رفتار^۲. برای بررسی‌ی تاثیر یادگیری در عمل‌کرد عامل دو معیار مختلف استفاده می‌شود. اولی میزان پاداش (و البته تنبیه) دریافتی عامل است که عمل‌کرد رفتاری عامل در طول یادگیری را نشان می‌دهد و دیگری احتمال خطا (۳-۶) است. معیار دوم نشان دهنده‌ی میزان دانش کسب شده‌ی عامل بدون توجه به استراتژی‌ی کاوشی به کار رفته است. پاداش کسب شده در شکل ۴-۶ و احتمال خطا در شکل ۵-۶ نمایش داده شده است. مطابق انتظار، عامل‌ای که از دانش اولیه‌ی ساختار مناسب برخوردار بوده است سریع‌تر از عامل با دانش اولیه‌ی نامناسب عمل‌کرد صحیح را یاد گرفته است. دلیل این موضوع همان‌طور که توضیح داده شد نیاز به یادگیری میزان کم‌تری از نگاشت‌های حالت-عمل است. به عبارت دیگر تنها کافی است که نگاشت حالت-عمل عامل‌ای که می‌تواند عمل مناسب را پیش‌نهاد کند یاد گرفته شود و دیگر نیازی به یادگیری NA در رفتارهای بالایی

^۲ برای بخش یادگیری ساختار از روش یادگیری ساختار با بازنمایی مرتبه صفر استفاده شده است.

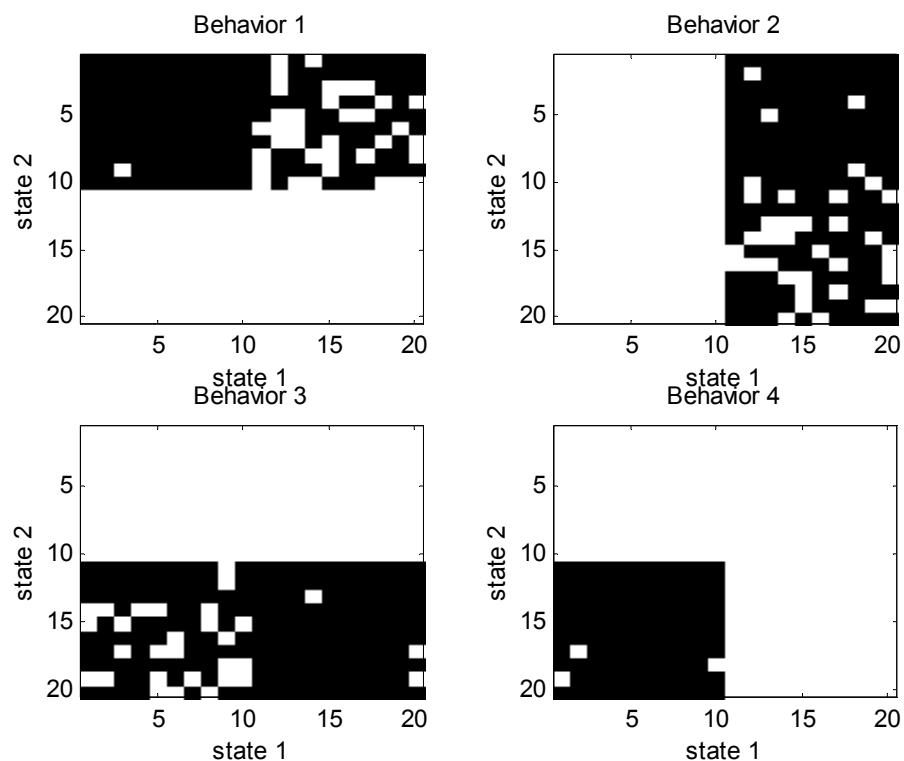
تحریک شده - در صورت وجود - نیست. عمل کرد یادگیری هم‌زمان ساختار و رفتار به دلیل آن که از اطلاعات اولیه‌ی کم‌تری نسبت به حالت یادگیری رفتار با ساختار مناسب استفاده می‌کند اندکی پایین‌تر است اما از حالت یادگیری رفتار با ساختار نامناسب به‌تر است. در این حالت یادگیری ساختار به کمک یادگیری رفتار آمده است و عامل به سمت ساختارهایی حرکت کرده که بیش‌تر پاداش‌آور بوده‌اند. کاوش موجود باعث پایین آمدن کیفیت عمل کرد سیستم شده است. میزان حداکثر پاداش قابل دریافت با آن سطح کاوش در شکل ۴-۶ نمایش داده شده است. دیده می‌شود که نتایج روش‌های پیش‌نهادی، به طور مجانبی به آن حد نزدیک می‌شوند. همچنین، بررسی‌ی احتمال خطا (۳-۶) مشخص می‌کند که در انتخاب اپیزود ۱۰۰۰۰ام (آخرین اپیزود)، احتمال خطا برای حالت‌های یادگیری رفتار (ساختار نادرست)، یادگیری هم‌زمان ساختار و رفتار و یادگیری رفتار (ساختار درست) به ترتیب برابر با 0.024، 0.009 و 0 است. نتایج حاصل بیان‌گر قابلیت روش‌های یادگیری پیش‌نهادی هستند. برای درک بیش‌تر چگونگی روند یادگیری، نواحی یادگرفته شده را در طول فرآیند یادگیری مشخص کرده‌ام (شکل‌های ۶-۶ تا ۶-۹). این نتایج حاصل یادگیری هم‌زمان رفتار و ساختار است. مطابق انتظار نواحی یادگرفته شده به تدریج بزرگ می‌شوند. جزییات این آزمایش در جدول ۱-۶ آمده است.



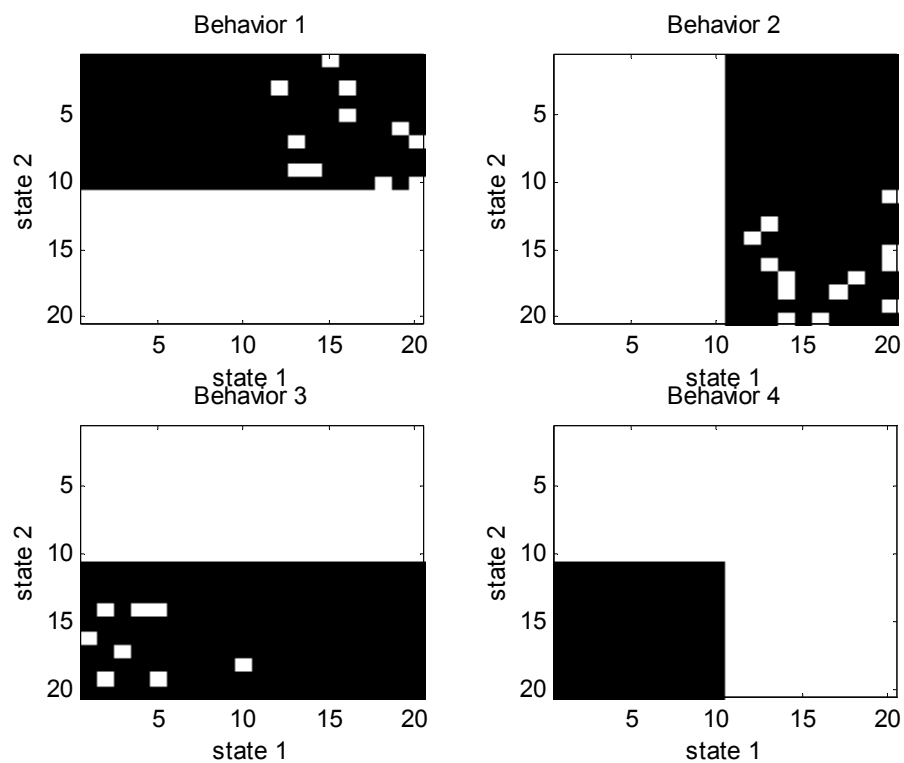
شکل ۶-۶. (مساله مجرد) نواحی یادگرفته شده (اپیزود ۱۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.



شکل ۶-۷. (مساله مجرد) نواحی یادگرفته شده (اپیزود ۵۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.



شکل ۸-۶. (مساله مجرد) نواحی یادگرفته شده (اپیزود ۱۵۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

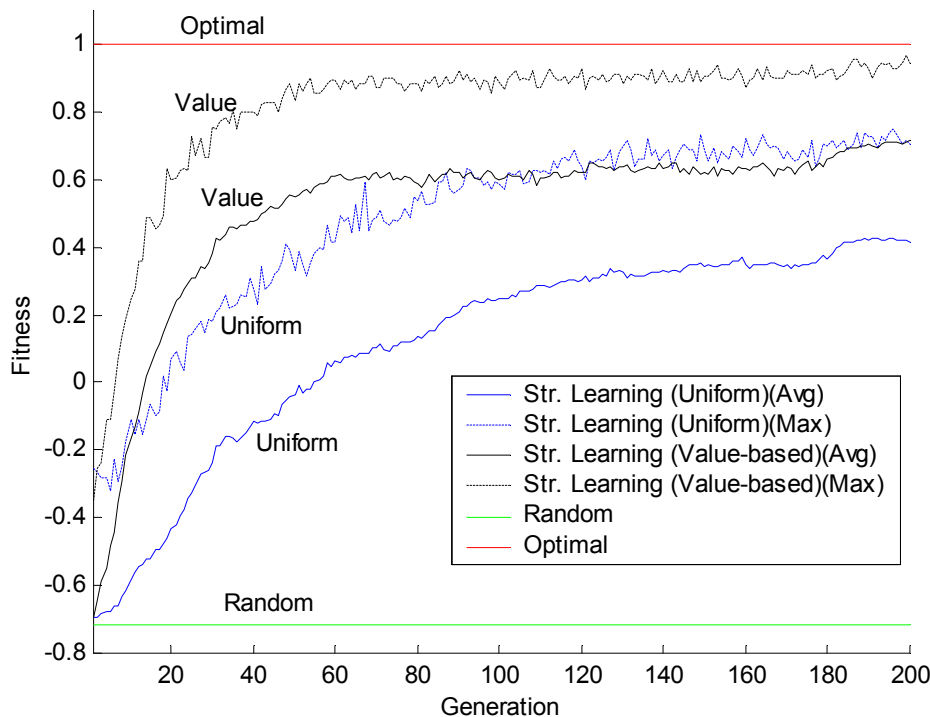


شکل ۸-۹. (مساله مجرد) نواحی یادگرفته شده (اپیزود ۳۰۰۰). نواحی تیره به درستی یاد گرفته شده‌اند.

۶-۱-۳) تکامل رفتار و یادگیری ساختار

در این بخش، ایده‌های مطرح شده در فصل ۵ را بر روی مساله‌ی مجرد آزمایش می‌کنم. بررسی خواهم کرد که آیا تکامل رفتارها می‌تواند جای‌گزین مناسب‌ای برای یادگیری باشد یا خیر و همچنین تاثیر ترکیب هم‌زمان یادگیری و تکامل را نیز مطالعه می‌کنم. تکامل رفتار را در دو حالت مختلف اشتراک سازگاری‌ی یک‌نواخت و اشتراک سازگاری‌ی ارزش‌محور بررسی می‌کنم. همچنین، تاثیر استفاده از ایده‌های الگوریتم ممیتیک - که در آن مجموعه‌ای از الگوهای فرهنگی شکل گرفته و اعضای جدید جامعه از ابتدای کنش‌شان با محیط به آن‌ها مجهز هستند - نیز مطالعه می‌شود. در این مساله، سازگاری‌ی عامل معادل با متوسط پاداش و تنبیه‌ای است که در طول دوره‌ی زندگی‌اش دریافت می‌کند. از آن‌جا که در مساله‌ی مجرد، پاسخ صحیح پاداش ۱ و پاسخ نادرست تنبیه ۱- دارد، عدد سازگاری بین ۱ و ۱- خواهد بود.

در این بخش، آزمایش‌های انجام شده بر روی مساله‌ای با ابعاد 5×5 و دارای هفت عمل صورت می‌گیرد. در این مساله هفت رفتار $B_i (i=1, \dots, 7)$ داریم که هر کدام از آن‌ها می‌توانند در کل ۲۵ حالت مساله تحریک شوند. خروجی هر رفتار، متناسب با شماره‌ی آن رفتار است: یعنی رفتار B_i می‌تواند یا عمل a_i را انتخاب کند و یا NA را. واضح است که برای حل مساله‌ی $\Pi(0)$ رفتارها می‌بایست یاد بگیرند که در چه شرایط لازم است فعال شوند و در چه شرایطی سکوت اختیار کنند. از آن‌جا که بازنمایی همه‌ی رفتارها یک‌سان است و فضای تحریک آن‌ها نیز تفاوت‌ای با هم ندارد، پس به ازای هر ترتیب رفتاری، نگاشت $A'_i : S'_i \rightarrow B_i(s')$ وجود دارد که به پاسخ مناسب می‌انجامد. البته توجه کنید که به ازای ساختارهای مختلف، این نگاشت یک‌سان نیست. توضیحات بیش‌تر در مورد اندازه‌ی جمعیت تکامل‌یابنده، پارامترهای یادگیری و پارامترهای الگوریتم ممیتیک در جدول ۲-۶ آمده است. اینک نتایج آزمایش را توضیح می‌دهم.



شکل ۶-۱۰. (مساله مجرد) مقایسه‌ی بین متوسط و حداکثر سازگاری جمعیت به ازای روش‌های مختلف اشتراک سازگاری همراه با یادگیری ساختار. اشتراک سازگاری یک‌نواخت (آبی)، اشتراک سازگاری ارزش‌محور (سیاه). خط سبز متوسط سازگاری معماری‌ای با ساختار و رفتار کاملاً تصادفی است (با احتمال انتخاب عمل و NA برابر) و خط قرمز حداکثر سازگاری دست‌یافتنی است. خطوط پرننگ متوسط سازگاری جمعیت را نشان می‌دهند و خطوط نقطه‌چین، حداکثر سازگاری را.

در شکل ۶-۱۰، مقایسه‌ای بین سازگاری روش اشتراک سازگاری یک‌نواخت و ارزش‌محور با پاسخ بهینه و پاسخ کاملاً تصادفی نشان شده است. دیده می‌شود که متوسط سازگاری جمعیت‌ای که از اشتراک سازگاری ارزش‌محور استفاده کرده است به طور قابل ملاحظه‌ای بالاتر از جمعیت‌ای است که از اشتراک سازگاری یک‌نواخت استفاده کرده. به‌ترین جمعیت روش اول نیز در حدود پاسخ بهینه - که در همه‌ی شرایط پاسخ صحیح تولید می‌کند و در نتیجه دارای سازگاری ۱ است - رفتار می‌کند. به‌ترین جمعیت با اشتراک سازگاری یک‌نواخت در حدود متوسط سازگاری جمعیت با سازگاری ارزش‌محور است. سازگاری پاسخ کاملاً تصادفی (که در آن هر رفتار با احتمال $\frac{1}{2}$ عمل خود را انتخاب کرده و با همان احتمال NA را انتخاب می‌کند) بسیار پایین‌تر از بقیه‌ی حالت‌ها است. نکته‌ی جالب و البته طبیعی این‌که مقدار سازگاری حالت تصادفی که در حدود $0.7 -$ است بسیار نزدیک به مقدار متوسط سازگاری در نسل اول روش‌های تکاملی است. برای محاسبه‌ی سازگاری حالت کاملاً

تصادفی دو روش به کار بردنی است: روش عددی و روش تحلیلی. نتیجه‌ی روش عددی را -که با اجرای چندین باره‌ی معماری‌ای تصادفی به دست آمده- در شکل می‌بینید. برای محاسبه‌ی تحلیلی این مقدار در ساختار مورد نظر و با فرض این‌که عمل‌های $\Pi(\cdot)$ به صورت یک‌نواخت انتخاب شده باشند، یعنی:

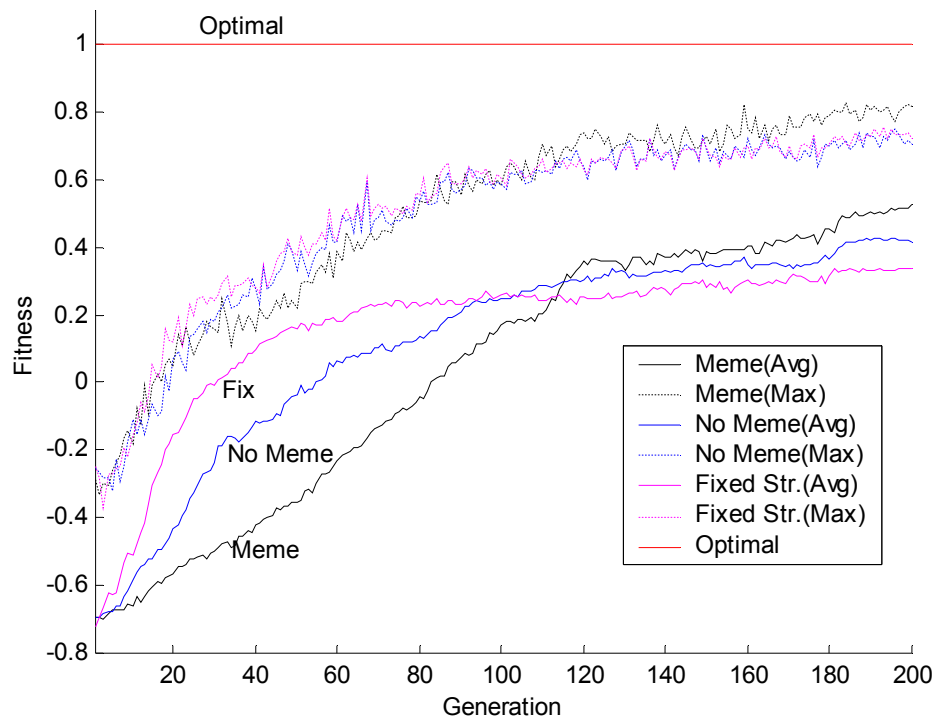
$$P(a_i \in \Pi) = P(a_j \in \Pi) \quad \forall a_i, a_j \in \Pi \quad (۵-۶)$$

امید ریاضی سیگنال تقویتی‌ی دریافتی را می‌توان بدین‌گونه محاسبه کرد:

$$E_{\pi_{random}}[r] = \frac{1}{m} \sum_{i=1}^m \left(\frac{1}{2^i} \times (1) + \left(1 - \frac{1}{2^i} \right) \times (-1) \right) = \frac{1}{m} \left[2 \sum_{i=1}^m \frac{1}{2^i} - m \right] \quad (۶-۶)$$

که در آن m مطابق معمول تعداد لایه‌ها (و رفتارها) است. جمله‌ی اول داخل اولین علامت جمع (با در نظر گرفتن $1/m$) بیان‌گر امید ریاضی‌ی پاداش‌ای است که از درست رفتار کردن لایه‌ی i ام حاصل می‌شود. برای درست عمل کردن آن لایه، لایه‌های بالا می‌بایست NA پیش‌نهاد دهند (هر کدام با احتمال $1/2$) و خود آن لایه نیز می‌بایست عمل درست را انتخاب کند (احتمال $1/2$). جمله‌ی دوم همان عبارت، میزان امید ریاضی‌ی دریافت تنبیه برای همان لایه است. نتیجه‌ی این مقدار برای معماری‌ی هفت لایه‌ی این بخش برابر با -0.7165 خواهد بود.

دلیل کیفیت برتر روش سازگاری‌ی ارزش‌محور، استفاده‌ی مناسب از دانش حاصل از یادگیری است. دانش حاصل از یادگیری که به صورت متوسطی از $V_{T(B_i)}$ ها منتقل می‌شود، اطلاعات دقیق‌ای در مورد عمل‌کرد هر رفتار دارند. در حالی که معیار سازگاری یک‌نواخت -که عمل‌کرد همه‌ی رفتارها با هم را در نظر می‌گیرد- به طور دقیق میزان سهم هم رفتار را مشخص نمی‌کند. علاوه بر این مورد، مساله‌ی مورد نظر دارای ویژگی‌ی جالبی است و آن هم این‌که بهینه‌سازی‌ی سازگاری‌ی هر رفتار به طور مستقیم به بهینه‌سازی‌ی سازگاری‌ی همه‌ی رفتارها می‌انجامد. دلیل این موضوع این است که سازگارتر

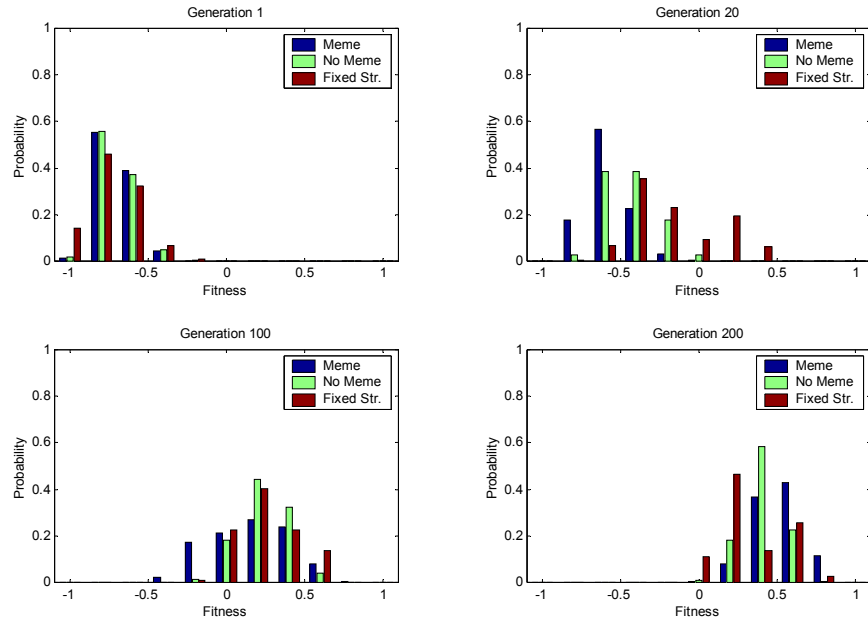


شکل ۱۱-۶. (مساله مجرد) مقایسه‌ی متوسط و حداکثر سازگاری به ازای روش‌های مختلفی که از اشتراک سازگاری ی یک‌نواخت استفاده می‌کنند: تکامل رفتارها و یادگیری ساختار بدون استفاده از مم (آبی)، تکامل رفتارها و یادگیری ساختار با استفاده از مم (سیاه)، تکامل رفتارها و ساختار ثابت (بنفش). خطوط پررنگ متوسط سازگاری جمعیت و خطوط نقطه‌چین حداکثر سازگاری را نشان می‌دهند. حداکثر سازگاری ممکن با خط قرمز نشان داده شده است.

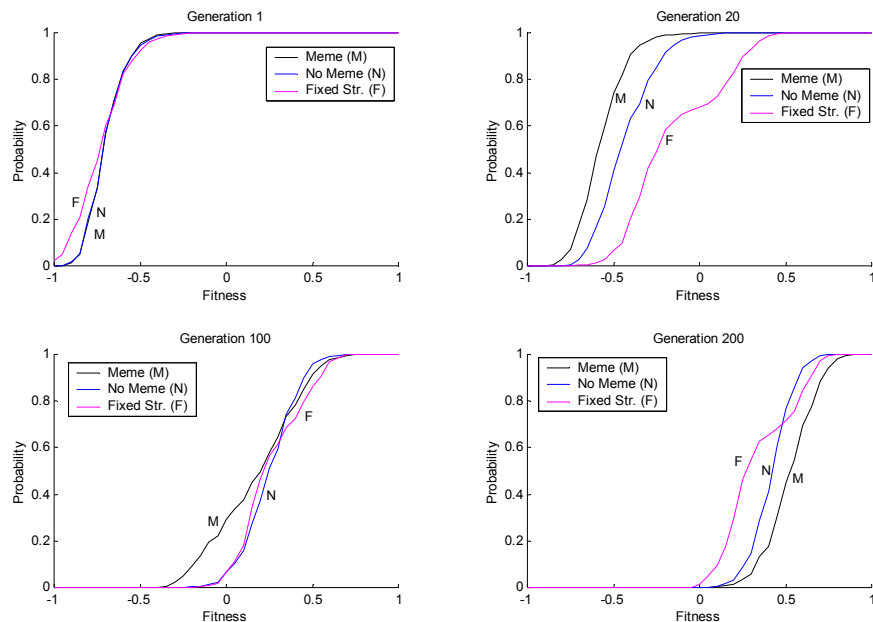
بودن هر رفتار - و بالا بودن امتیاز آن - معادل با سازگارتر بودن کل رفتارها است: ممکن نیست رفتاری سازگارتر باشد ولی کل رفتارها، معماری‌ای سازگارتر نسازند. این ویژگی در همه‌ی مسایل وجود ندارد. به همین دلایل است که در این مساله سرعت تکامل در روش اشتراک سازگاری ارزش‌محور به طور قابل ملاحظه‌ای بالاتر از سرعت تکامل در روش اشتراک سازگاری ی یک‌نواخت سریع‌تر است.

اینک تاثیر یادگیری بر تکامل و همچنین تاثیر وجود مم بر یادگیری و در نتیجه بر تکامل بررسی می‌شود. در شکل ۱۱-۶، مقایسه‌ای بین چندین حالت مختلفی که همه در بخش تکامل رفتارشان از اشتراک سازگاری ی یک‌نواخت سود می‌جویند انجام شده است. مطابق انتظار و تحلیل‌هایی که در بخش‌های مربوط به یادگیری نیز انجام شده است هرچه دانش اولیه‌ی بیشتری از مساله داشته باشیم، یادگیری به‌تر و سریع‌تر خواهد بود. برای همین انتظار می‌رود که اضافه کردن دانش ساختار مناسب باعث سرعت بخشیدن به فرآیند تکامل شود. چنین پدیده‌ای را در شکل ۱۱-۶ مشاهده می‌کنید.

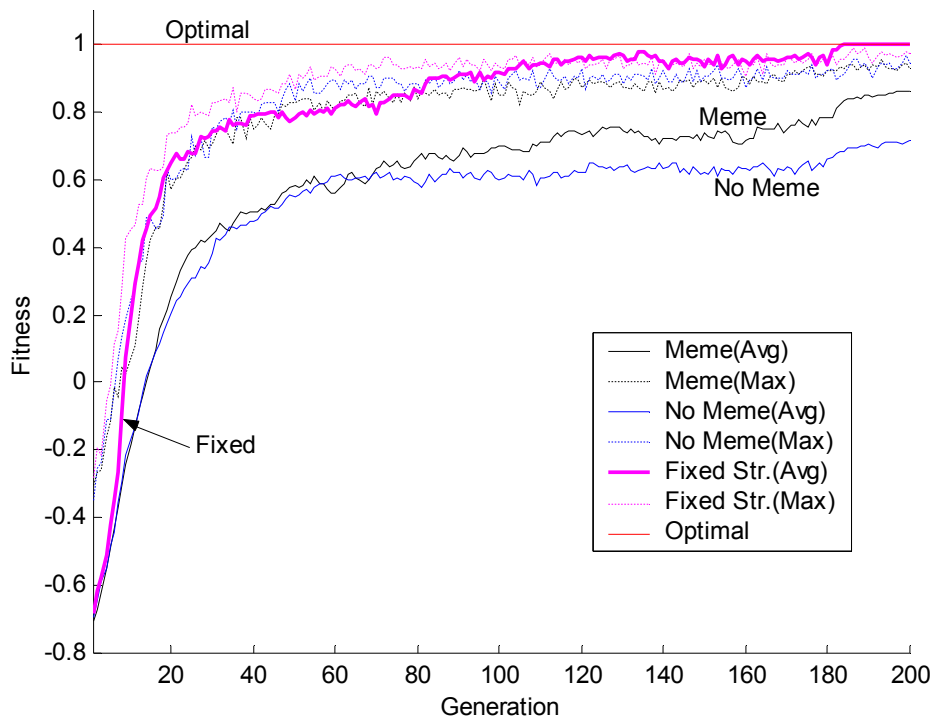
سرعت یادگیری حالت ساختار ثابت از دو حالت دیگری که از یادگیری استفاده می‌کنند بیش‌تر است. اما نکته‌ای که در این مساله مجرد وجود دارد و در ابتدای بخش توضیح دادم این است که این مساله – برخلاف مساله‌ی استفاده شده در آزمایش‌های مربوط به یادگیری – دارای پاسخ یگانه برای ساختار نیست: به ازای هر ترکیب ساختاری غیرتکراری، نگاشت متناسب حالت-عمل برای رفتارها وجود دارد. پس ثابت نگاه داشتن ساختار بیش از آن که عمل‌کردی ناشی از تزریق دانش صحیح داشته باشد، از پراکندگی ساختارهایی که عامل با آن‌ها کار می‌کند (و پراکندگی موجب پایین آمدن عمل‌کرد می‌شود) جلوگیری می‌کند. به عبارت دیگر اگر تعدادی رفتار برای ترتیب خاص T_1 تکامل یافته باشند، آن‌گاه در صورت‌ای که در آینده در ساختار T_2 به کار روند، بازده عامل پایین می‌آید. یادگیری برای پیدا کردن به‌ترین ترتیب کمی زمان لازم دارد و در نتیجه طبیعی است که کیفیت‌اش حداقل در چند اپیزود اولیه پایین باشد. حالت استفاده از یادگیری کندتر از ساختار ثابت است و حالت‌ای که یادگیری‌اش با استفاده از مم تقویت شده باشد نیز از هر دو کندتر است. اما در نسل‌های نهایی مشاهده می‌شود که حالت‌هایی که یادگیری دارند کمی به‌تر از حالت بدون یادگیری عمل می‌کنند. دلیل این موضوع "شاید" این باشد که عامل با استفاده از یادگیری توانسته است ترتیب‌ای از رفتارهای هنوز--به-طور-کامل تکامل نیافته‌ای را بیابد که باعث عمل‌کرد به‌ترش می‌شوند. در شکل ۱۲-۶، چگالی احتمال و در شکل ۱۳-۶، احتمال تجمعی شایستگی عامل‌ها در چند نسل مختلف مشخص شده است. هیستوگرام شکل ۱۲-۶ هر چقدر در سمت راست تجمع بیش‌تری داشته باشد، به معنای وجود عامل‌هایی با کیفیت بالاتر است. تمایل به سمت راست در شکل ۱۳-۶ نیز به همان معنا است. دیده می‌شود که در نسل ۱ که هنوز تکامل‌ای صورت نگرفته عامل‌هایی که یادگیری داشته‌اند اندکی از عامل‌های با معماری ثابت به‌تر بوده‌اند. این به آن علت است که یادگیری توانسته ترتیب‌ای از رفتارها را بیابد که به بالا رفتن سازگاری عامل منجر شود. این پدیده در شکل ۱۳-۴ واضح‌تر است. از نمودارها بر می‌آید که پس از گذشت چند نسل، نمودار به سمت راست تمایل پیدا می‌کند. در ابتدای تکامل (مثلاً نسل ۲۰)، تکامل با ساختار ثابت به وضوح جلوتر از تکامل با یادگیری است ولی پس از گذشت مدت زمان کافی برای تکامل کیفیت همه در حدود هم می‌شود و حتی در آخرین نسل (نسل ۲۰۰)، حالت‌های دارای یادگیری – و مخصوصاً حالت تقویت شده با الگوریتم ممتیک – جلو می‌افتد.



شکل ۱۲-۶. (مسالهی مجرد) مقایسه‌ی چگالی احتمال سازگاری برای تکامل با اشتراک سازگاری یک‌نواخت در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری اولیه ساختار به کمک مم (آبی تیره)، عامل‌هایی با ساختار اولیه تصادفی (سبز) و عامل‌هایی با ساختار ثابت (قرمز) انجام گرفته است.



شکل ۱۲-۶. (مسالهی مجرد) مقایسه‌ی توزیع احتمال تجمعی ($P\{Fitness \leq f\}$) سازگاری برای تکامل با اشتراک سازگاری یک‌نواخت در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری اولیه ساختار به کمک مم (سیاه)، عامل‌هایی با ساختار اولیه تصادفی (آبی) و عامل‌هایی با ساختار اولیه ثابت (بنفش) انجام گرفته است. تمایل نمودار به سمت راست نشان‌گر بالاتر بودن احتمال تولید عامل‌های سازگارتر است.



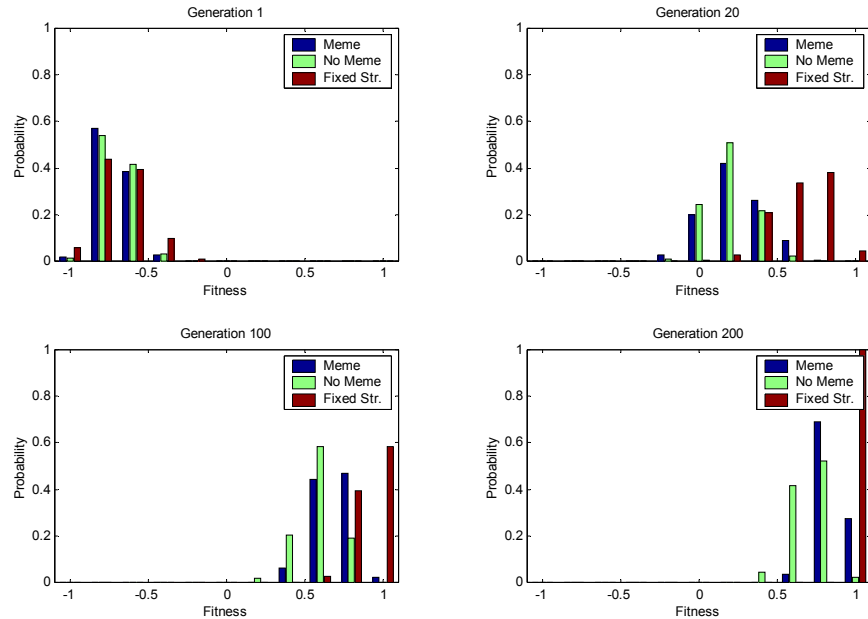
شکل ۶-۱۴. (مساله مجرد) مقایسه‌ی متوسط و حداکثر سازگاری به ازای روش‌های مختلفی که از اشتراک سازگاری ارزش‌محور استفاده می‌کنند: تکامل رفتارها و یادگیری ساختار بدون استفاده از مم (آبی)، تکامل رفتارها و یادگیری ساختار با استفاده از مم (سیاه)، تکامل رفتارها و ساختار ثابت (بنفش). خطوط پررنگ متوسط سازگاری جمعیت و خطوط نقطه‌چین حداکثر سازگاری را نشان می‌دهند. حداکثر سازگاری ممکن با خط قرمز نشان داده شده است.

مشابه آزمایش‌های ذکر شده برای تکامل با اشتراک سازگاری ارزش‌محور نیز انجام شده است.

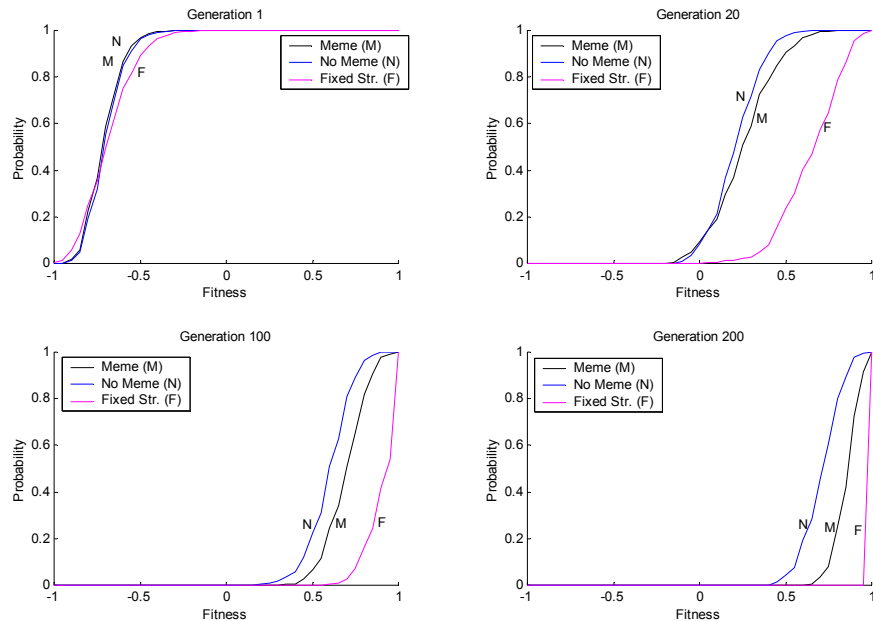
در شکل ۶-۱۴ می‌بینید که تکامل رفتار با ساختار ثابت بسیار سریع به نتایج قابل رقابت با حالت بهینه می‌رسد. حالت‌های مربوط به یادگیری از آن کندترند و متوسط سازگاری جمعیت آن‌ها پایین‌تر از حالت ساختار ثابت است. مطابق آن‌چه پیش‌تر گفته شد، چنین نتیجه‌ای قابل انتظار است. حالت‌ای که از الگوریتم مم‌تیک و انتقال فرهنگ استفاده می‌کند در نهایت به سازگاری متوسط بالاتری می‌رسد. دلیل این موضوع این است که انتقال ساختارهای مناسب با رفتارهای تکامل یافته، کار یادگیری هر عامل را بسیار ساده‌تر می‌کند و عامل از همان قدم اول -و یا چند قدم ابتدایی- به ساختار مناسب برای مجموعه رفتارهای در اختیارش دست می‌یازد. به‌ترین عامل‌های ایجاد شده البته در حد و حدود هم‌اند و بسیار نزدیک به حالت بهینه. هیستوگرام (شکل ۶-۱۵) و توزیع تجمعی سازگاری (شکل ۶-۱۶) نیز اطلاعات مفیدی به ما می‌دهند. به عنوان نمونه در نسل ۲۰ دیده می‌شود که حالت ثابت به عامل‌هایی

می‌انجامد که کیفیت بسیار بالاتری نسبت به همتهای یادگیر خود دارند. این اختلاف البته هم‌چنان در نسل‌های بعد نیز وجود دارد اما فاصله به تدریج کم می‌شود (شکل ۱۶-۶ را ببینید). هم‌چنین اختلاف حالت دارای مم و حالت بدون آن در نسل‌های بالا - که رفتارها تقریباً شکل گرفته‌اند - قابل ملاحظه است. همان‌طور که انتظار می‌رفت وجود مم باعث به‌تر شدن کیفیت کار و سریع‌تر شدن یادگیری می‌شود. وجود مم به نوعی بعضی از ویژگی‌های مثبت ساختار ثابت را به فرآیند تولید عامل اضافه می‌کند در حالی که نیازی به دانستن از-پیش ساختار ندارد.

در این آزمایش‌ها دیده شد که تکامل رفتار می‌تواند جای‌گزینی برای یادگیری رفتار باشد. البته به دلیل نداشتن ارتباط مستقیم با محیط و تنها قابلیت مشاهده‌ی نتیجه‌ی کلی‌ی عمل‌کرد عامل، سرعت تکامل و میزان محاسبه‌ی لازم برای آن بسیار کم‌تر از یادگیری رفتار است. مخصوصاً به این نکته توجه کنید که فضای مساله در این حالت بسیار کوچک‌تر از فضای مساله‌ی انجام شده برای آزمون یادگیری بوده است و عملاً با مساله‌ای ساده‌تر طرف بودیم. با این حال در صورت‌ای که بتوان پاسخ بهینه‌ی مساله را به پاسخ بهینه‌ی تعدادی زیرمساله (ارزش هر رفتار) تجزیه کرد (مانند مساله‌ی مورد بررسی)، روش اشتراک سازگاری ارزش‌محور باعث سرعت بخشیدن به تکامل می‌شود.



شکل ۱۵-۶. (مساله‌ی مجرد) مقایسه‌ی چگالی احتمال سازگاری برای تکامل با اشتراک سازگاری ارزش‌محور در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری‌ی اولیه ساختار به کمک مم (آبی تیره)، عامل‌هایی با ساختار اولیه‌ی تصادفی (سبز) و عامل‌هایی با ساختار ثابت (قرمز) انجام گرفته است.



شکل ۱۶-۶. (مساله‌ی مجرد) مقایسه‌ی توزیع احتمال تجمعی ($P\{Fitness \leq f\}$) سازگاری برای تکامل با اشتراک سازگاری ارزش‌محور در چندین نسل مختلف. مقایسه بین عامل‌هایی با سوگیری‌ی اولیه ساختار به کمک مم (سیاه)، عامل‌هایی با ساختار اولیه تصادفی (آبی) و عامل‌هایی با ساختار اولیه ثابت (بنفش) انجام گرفته است. تمایل نمودار به سمت راست نشان‌گر بالاتر بودن احتمال تولید عامل‌های سازگارتر است.



شکل ۱۷-۶. سه روبات در حال بلند کردن جسم‌ای بزرگ

۶-۲) بلند کردن جسم چند-روباتی

با وجود آن‌که مساله‌ی مجرد سنگ محک مناسب و کافی‌ای برای ارزیابی‌ی روش‌های پیش‌نهادی است، اما حل مساله‌ای از دنیای واقعی نیز قابل توجه خواهد بود. مساله‌ی بلند کردن جسم توسط چند روبات را به عنوان مساله‌ی دیگر این پژوهش انتخاب کردم. این مساله پیش‌تر توسط معماری رفتارگرایی سلسله‌مراتبی‌ای مشابه با ساختار PPSSA به کار رفته در این پژوهش توسط طراح و با روش سعی و خطا حل شده بود [Nili01]. در آن پژوهش ساختار و هر کدام از رفتارهای معماری با سعی و خطای بسیار پیدا شده بودند. از آن‌جا که هدف اصلی پژوهش فعلی یافتن روش‌هایی برای خودکار کردن فرآیند طراحی است، تلاش برای حل این مساله -که طراحی‌ای وقت‌گیر و دشوار داشته است- جالب توجه می‌نماید. در صورتی که روش‌های این پایان‌نامه بتواند به بازده‌ای قابل مقایسه با طراحی دستی -که حاصل هوش انسانی است- برسد، آن‌گاه می‌توان ادعا کرد که روش‌های پیش‌نهادی قابل

اعتنا هستند. پیش از توضیح بیش‌تر درباره‌ی روش کار، بیش‌تر راجع به مساله‌ی بلند کردن یک جسم توسط چند روبات توضیح می‌دهم.

فرض کنید که قرار باشد گروه‌ای از روبات‌ها جسم‌ای بزرگ و نامشخص را بلند کنند (شکل ۱۷-۶). جسم به اندازه‌ای است که هیچ کدام از روبات‌ها به تنهایی نمی‌توانند آن را در بر گرفته و بلند کنند و در نتیجه هم‌کاری‌ی چند روبات با هم ضروری می‌نماید. در این شرایط احتمالا مناسب‌ترین روش بلند کردن جسم استفاده از مکانیزم چنگالی^۳ باشد. در هنگام بلند کردن جسم با توجه به موقعیت نسبی‌ی هر روبات به مرکز جرم جسم، نیروی اعمالی‌ی لازم هر روبات متفاوت با دیگری خواهد بود و این منجر به ناهمگنی بین رفتار روبات‌های تیم می‌شود. هر روبات باید قابلیت کار در گستره‌ای از بارگذاری‌های مختلف را داشته باشد. هم‌چنین در هنگام بلند کردن جسم هر کدام از روبات‌ها می‌باید به گونه‌ای حرکت کند تا لغزشی در نقطه‌ی اتصال‌شان با جسم به وجود نیاید. اما اگر بتوان زاویه‌ی کج‌شده‌گی^۴ جسم را تا حد ممکن کوچک نگه داشت، با اضافه کردن گیره‌ای^۵ در نقطه‌ی تماس موازی با سطح زیرین جسم، دیگر نیازی به حرکت کردن روبات‌ها وجود ندارد. هم‌چنین کم نگه داشته زاویه‌ی کج‌شده‌گی جسم از برخورد جسم با روبات‌ها جلوگیری می‌کند.

برای پایدار نگاه داشتن جسم در هنگام بلندکردن یا حرکت کافی است که جسم به گونه‌ای حرکت کند که Zero Moment Point (ZMP) آن درون سطح‌ای بیافتد که راس‌های آن سطح از نقطه‌ی تماس جسم و روبات تشکیل شده‌اند. مثلا اگر سه روبات داریم (مانند مساله‌ی فعلی) ZMP آن در درون مثلث نقطه‌ی تماس جسم و روبات بیافتد. حال با توجه شکل جسم و مکان مرکز جرم جسم می‌توان زاویه‌ی کج‌شده‌گی‌ای را به دست آورد که ZMP آن به مرز ناحیه‌ی ذکر شده برسد. پس اگر بتواند زاویه‌ی کج‌شده‌گی را همیشه کوچک‌تر از آن مقدار مشخص شده نگاه داشت، روبات‌ها می‌توانند بدون هیچ حرکت افقی‌ای جسم را بلند کنند، جسم به روبات‌ها برخورد نخواهد کرد و سیستم نیز پایدار خواهد بود. برای حل این مساله نیازی به کنترل‌کننده‌ی مرکزی نیست. نشان داده شده است که هر روبات با اطلاعات محلی‌ای که به دست می‌آورد می‌تواند این مساله را حل کند.

³ Fork Lifting Mechanism

⁴ Tilt Angle

⁵ Compliance

به طور خلاصه، مساله مورد نظر، چگونگی بلند کردن هم‌کارانه‌ی جسم‌ای نامشخص به ارتفاع‌ای از پیش معلوم و در همان زمان کوچک نگه داشتن زاویه‌ی کج‌شده‌گی، توسط سه روبات یک‌سان است. در این مساله کنترل مرکزی یا ارتباطی بین روبات‌ها وجود ندارد و روبات‌ها تنها از اطلاعات محلی‌شان استفاده می‌کنند. این مساله را هم با استفاده از یادگیری و هم با استفاده از تکامل حل می‌کنم. روش‌های یادگیری را در سه مرحله‌ی مختلف بررسی می‌کنم: یافتن ساختار مناسب برای روبات‌ها با فرض داشتن رفتارهای مناسب، یادگیری نگاشت مناسب برای هر رفتار با فرض دانستن ساختار صحیح و یادگیری هم‌زمان ساختار و رفتار. ترکیب تکامل و یادگیری نیز در این مساله در دو حالت تکامل رفتار با فرض دانستن ساختار و تکامل رفتار و یادگیری ساختار توأمان انجام می‌شود. آزمایش‌های مختلفی به ازای روش‌های گوناگون اشتراک سازگاری (یک‌نواخت و ارزش-محور) انجام شده و هم‌چنین تاثیر استفاده از مم بر روی بازده سیستم بررسی می‌شود.

۶-۲-۱) یادگیری ساختار

در گام اول فرض می‌کنم که مجموعه‌ای از رفتارهای مناسب در اختیار عامل قرار گرفته است و عامل می‌بایست رفتارها را به گونه‌ای بچیند تا وظیفه مورد نظر به درستی انجام شود. فرض داشتن مجموعه‌ای از رفتارهای مناسب در بعضی از موارد عملی قابل قبول است چون گاهی طراح درک نسبی‌ای از مساله‌ی مورد نظر و رفتارهای مناسب برای آن دارد ولی آگاهی‌ای نسبت به اولویت نسبی‌ی آن رفتارها و در نتیجه موقعیت‌شان در معماری ندارد. با این همه این شرط را نیز در گام‌های بعد بر می‌دارم و به عامل اجازه می‌دهم تا رفتارها را نیز یاد بگیرد.

این مساله را با دو روش یادگیری ساختار با بازنمایی مرتبه صفر و مرتبه یک حل می‌کنم. در این پژوهش از مجموعه رفتارهای مشابه با [Nili01] استفاده شده است. اگر $z(k)$ را ارتفاع نقطه‌ی تماس روبات-جسم، $v(k)$ سرعت بالا رفتن آن نقطه و $\tau(k)$ به عنوان زاویه‌ی کج‌شده‌گی روبات در گام k ام بگیریم^۶، رفتارها بدین گونه توصیف می‌شوند:

^۶ توجه کنید که همه‌ی این مقادیر به صورت محلی قابل اندازه‌گیری‌اند.

$$\text{Push more: } v(k+1) = v(k) + \Delta v \quad (6-7)$$

$$\text{Do not go fast: if } v(k) > v_{\max} \text{ then } v(k+1) = v_{\max} \text{ else do nothing} \quad (6-8)$$

$$\text{Stop at goal: if } z(k) \geq z_{\text{goal}} \text{ then stop } (v(k+1) = 0) \quad (6-9)$$

$$\text{Hurry up: if } \tau(k) > \tau_0 \text{ and the robot is the lowest one then} \\ v(k+1) = \min(v(k) + \Delta v, v_{\max}) \quad (6-10)$$

$$\text{Slow down: if } \tau(k) > \tau_0 \text{ and the robot is the highest one then} \\ v(k+1) = \max(v(k) - \Delta v, 0) \text{ and if it is the middle one } v(k) = \max(v(k) - \Delta v/2, 0) \quad (6-11)$$

در شبیه‌سازی آزمایش‌های انجام گرفته برای مساله‌ی بالا بردن جسم توسط چند روبات از پارامترهای $\tau_0 = 5^\circ$ و $z_{\text{goal}} = 3$ ، $v_{\max} = 5$ ، $\Delta v = 1$ و زمانی $\Delta T = 0.005$ بوده است.

یکی از مسایل مهم در یادگیری تقویتی طراحی سیگنال تقویت مناسب است ([Dorigo94]، [Dorigo97] و [Ng99]). فصل ۸ این پایان‌نامه را ببینید. عمل کرد قابل قبول عامل به شرطی حاصل می‌شود که همه‌ی آنچه مورد نظر طراح بوده است به نوعی در سیگنال تقویت گنجانده شده باشد. اما این کار مخصوصاً به هنگامی که مساله‌ی مورد نظر کمی پیچیده باشد و مثلاً بیش از یک هدف در آن مورد نظر باشد کاملاً دشوار می‌شود. تاکنون راه کلی‌ای برای طراحی آسان این سیگنال به وجود نیامده است و در اکثر موارد با استفاده از روش سعی و خطا سعی در پیدا کردن سیگنال تقویت مناسب می‌کنند. برای مساله‌ی یادگیری ساختار سیگنال تقویت زیر مناسب تشخیص داده شد^۷:

^۷ توجه کنید که سیگنال تقویت به روبات فعلی اعمال می‌شود. در نتیجه مقادیری چون $v(k)$ و $z(k)$ می‌بایست به صورت مقادیر مکان و سرعت نقطه‌ی تماس روبات فعلی با جسم تفسیر شوند.

$$r_k = r_k^{\tau_1} + r_k^{\tau_2} + r_k^{z_1} + r_k^{z_2} + r_k^v \quad (6-12)$$

$$r_k^{\tau_1} = \begin{cases} 1 & \tau(k) - \tau(k-1) < -0.5 \\ -0.1 & \text{otherwise} \end{cases} \quad (6-13)$$

$$r_k^{\tau_2} = \begin{cases} \frac{1}{\sqrt{k}} & \tau(k) < \tau_0 \\ -0.1\sqrt{k} & \text{otherwise} \end{cases} \quad (6-14)$$

$$r_k^{z_1} = \begin{cases} 1 & |z(k) - z_{goal}| < 0.5 \\ -0.1 & \text{otherwise} \end{cases} \quad (6-15)$$

$$r_k^{z_2} = -1 \quad z(k) > z_{goal} + 0.2 \quad (6-16)$$

$$r_k^v = -0.1 \quad v(k) > v_{max} \quad (6-17)$$

رابطه‌ی (۶-۱۲)، سیگنال تقویت‌ای است که عامل دریافت می‌کند. این سیگنال از چندین بخش مختلف تشکیل شده است که هر کدام برقراری بخش‌ای از معیارهای طراحی را تقویت می‌کنند. همان‌طور که دیده می‌شود، (۶-۱۳) کم شدن زاویه‌ی کج‌شده‌گی را پاداش داده و زیاد شدن‌اش را تنبیه می‌کند. رابطه‌ی (۶-۱۴) کوچک بودن زاویه‌ی کج‌شده‌گی را پاداش داده و زیاد بودن‌اش را تنبیه می‌کند. این بخش از سیگنال تقویت کم بودن زاویه در شروع بلند کردن را بیش از کم بودن آن در انتهاب اپیزود تقویت می‌کند و به طور عکس، زیاد بودن زاویه را در انتهای اپیزود بیش از زیاد بودن آن در ابتدای مسیر تنبیه می‌کند. با این نوع تقویت، عامل یاد می‌گیرد تا هر چه سریع‌تر زاویه کج‌شده‌گی را کاهش دهد. رابطه‌ی (۶-۱۵) نزدیک هدف بودن را پاداش داده و دور بودن از آن را تنبیه می‌کند و (۶-۱۶) نیز عبور از هدف و نایستادن در نزدیکی آن را تنبیه می‌کند. در نهایت (۶-۱۷) حرکت بسیار سریع روبات را تنبیه می‌کند.

برای محک زدن کیفیت عمل‌کرد سیستم، دو معیار مختلف طراحی شده است. اولین آن‌ها -به طور تقریباً بدیهی و مرسوم- متوسط سیگنال تقویتی (r_k) دریافتی توسط عامل است. این مقدار نشان می‌دهد که یادگیری تا چه حدی موفق بوده است. اما توجه به این نکته مهم است که حتی اگر سیگنال تقویت نادرست طراحی شده باشد این مقدار به خاطر عمل یادگیری بیشینه می‌شود در حالی که ممکن است بیشینگی آن معادل با عمل‌کرد مطلوب -از دید طراح- نباشد. به همین دلیل، معیار

رفتارگرایانه‌ای طراحی شده است که نمایان‌گر رفتار مطلوب عامل است. این معیار با توجه به آنچه از رفتار روبات بالا برنده‌ی "خوب" انتظار داریم طراحی شده است:

$$\text{Performance} = \text{Goal Closeness} + \frac{1}{3} \text{Tilt Angle Stabilization} \quad (۶-۱۸)$$

$$\text{Goal Closeness} = \frac{1}{N} \sum_{k=1}^N w_k \left[|z_k - z_{\text{goal}}| < 0.5 \right] \quad (۶-۱۹)$$

$$\text{Tilt Angle Stabilization} = \frac{1}{N} \sum_{k=1}^N w_k \left[\tau(k) < 10^\circ \right] \quad (۶-۲۰)$$

که در آن

$$w_k = \begin{cases} 1 & 0 \leq k \leq \frac{2}{3}N \\ 2 & \frac{2}{3}N < k \leq N \end{cases} \quad (۶-۲۱)$$

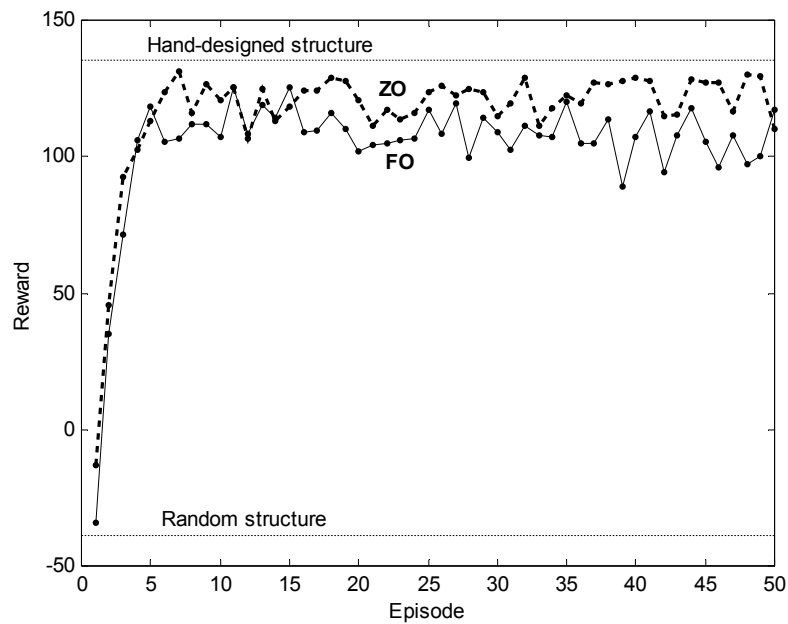
و N نیز طول اپیزود است. این معیار نزدیکی به هدف و کم بودن زاویه‌ی کج‌شده‌گی در انتهای مسیر را مطلوب می‌دارد.

در شکل ۶-۱۸، متوسط پاداش کسب شده‌ی عامل یادگیر با بازنمایی مرتبه صفر و مرتبه یک در طول ۵۰ اپیزود با ساختار طراحی شده توسط انسان^۸ و همچنین ساختار تصادفی مقایسه شده است. مقایسه‌ای مشابه با متوسط معیار رفتارگرایانه در شکل ۶-۱۹ نشان داده شده است. از آن‌جا که تعداد اپیزودها برخلاف مساله‌ی مجرد بسیار کم است و همچنین هر اپیزود در این مساله، تجربه‌ای کامل است (سعی در بلند کردن جسم از ابتدا تا موفقیت یا شکست) از هیچ فیلتر متوسط‌گیر لغزان^۹ استفاده نشده است. مشاهده می‌شود که هر دو روش یادگیری مبتنی بر دو بازنمایی مختلف عمل‌کرد مطلوبی داشته‌اند و روبات‌ها در اپیزودهای اندکی (به طور متوسط کم‌تر از ۱۰ اپیزود) توانسته‌اند به راه حل مساله دست

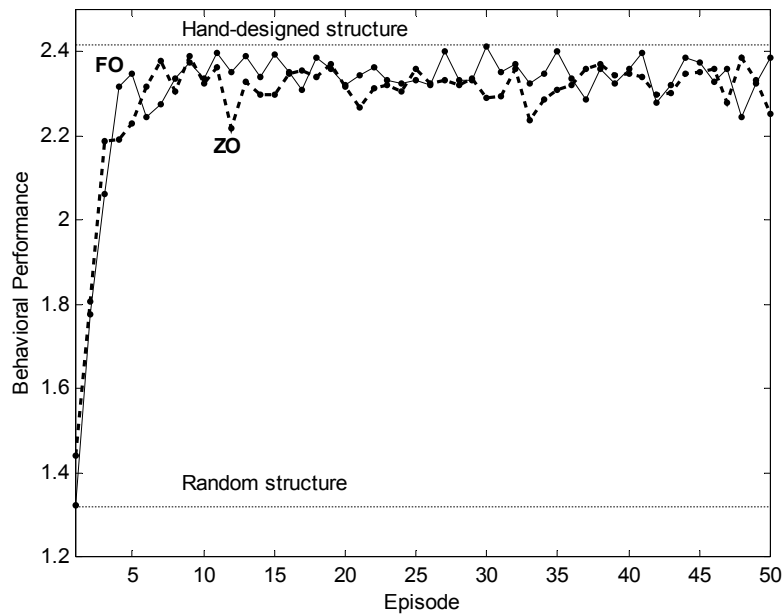
^۸ ساختار طراحی شده توسط انسان این بوده است: $T = [B_{\text{Push more}} \quad B_{\text{Don't go fast}} \quad B_{\text{Hurry up}} \quad B_{\text{Slow down}} \quad B_{\text{Stop}}]$

^۹ Moving Average

یابند. کیفیت کار هر دو بازنمایی قابل قبول است و هیچ کدام به طور مشخص‌ای برتر از دیگری نیست. همچنین سطح متوسط مقادیری که ساختارهای تصادفی کسب می‌کنند به وضوح پایین‌تر از سطح عامل‌های یادگیرنده است و می‌توان مطمئن بود که نتایج حاصل از به خاطر قابلیت یادگیری بوده است و کیفیت اتفاقی نتایج. همچنین مشاهده می‌شود که هر دوی معیارهای متوسط سیگنال تقویت دریافتی و معیار عمل‌کرد رفتارگرایانه در طول یادگیری افزایش می‌یابند. بررسی دقیق‌تر نشان می‌دهد که این دو دارای ضریب همبستگی 0.99 هستند. این نشان می‌دهد که بهینه‌سازی سیگنال تقویت طراحی شده باعث بهینه‌گی معیار رفتارگرایانه می‌شود و در نتیجه سیگنال تقویت طراحی شده تا حد قابل قبولی معادل همان معیار عمل‌کرد رفتارگرایانه مطلوب است. به مانند قبل جزییات شبیه‌سازی را در جدول ۶-۱ می‌توانید مشاهده کنید.



شکل ۱۸-۶. (مساله‌ی بلندکردن جسم) مقایسه متوسط سیگنال تقویتی دریافتی بین روش‌های یادگیری ساختار مبتنی بر بازنمایی مرتبه صفر (ZO) و بازنمایی مرتبه یک (FO)، ساختار از پیش طراحی شده و ساختار اتفاقی. رفتارها توسط انسان طراحی شده‌اند.



شکل ۱۹-۶. (مساله‌ی بلندکردن جسم) مقایسه معیار رفتارگرایانه بین روش‌های یادگیری ساختار مبتنی بر بازنمایی مرتبه صفر (ZO) و بازنمایی مرتبه یک (FO)، ساختار از پیش طراحی شده و ساختار اتفاقی. رفتارها توسط انسان طراحی شده‌اند.

۶-۲-۲) یادگیری رفتار و یادگیری هم‌زمان رفتار/ساختار

در گام بعد فرض وجود جعبه‌ابزاری از رفتارها را به کنار می‌نهیم و اجازه می‌دهیم عامل خود نگاشت مناسب حالت-عمل هر رفتار را یاد بگیرد. در این گام ساختار از پیش مشخص است. فرض این مرحله دانستن فضای ورودی و خروجی هر رفتار است. آن فضاها را بدین‌گونه تعریف می‌کنم:

$$\text{Push more: } S'_{\text{Push more}} = \{\emptyset\}, A'_{\text{Push more}} = \{v(k+1) = 1, NA\} \quad (۶-۲۲)$$

$$\text{Do not go fast}^{10}: S'_{\text{Don't go fast}} = \{v(k) | v(k) \in \{0^-, 1, 2, 3, 4, 5^+\}\},$$

$$A'_{\text{Don't go fast}} = \{v(k+1) = \max(0, v(k) - 0.5), NA\} \quad (۶-۲۳)$$

$$\text{Stop at goal: } S'_{\text{Stop}} = \{z(k) - z_{\text{goal}} < -0.2, z(k) - z_{\text{goal}} \geq -0.2\},$$

$$A'_{\text{Stop}} = \{v(k+1) = 0, NA\} \quad (۶-۲۴)$$

Hurry up:

$$S'_{\text{Hurry}} = \{\text{lowest robot, middle robot, highest robot}\} \times \{\tau < \tau_0, \tau > \tau_0\},$$

$$A'_{\text{Hurry}} = \{v(k+1) = 1, NA\} \quad (۶-۲۵)$$

Slow down:

$$S'_{\text{Slow}} = \{\text{lowest robot, middle robot, highest robot}\} \times \{\tau < \tau_0, \tau > \tau_0\},$$

$$A'_{\text{Slow}} = \{v(k+1) = -1, NA\} \quad (۶-۲۶)$$

و به عنوان سیگنال تقویت نیز از تابع زیر استفاده می‌کنم:

$$r_k = r_k^\tau + r_k^{z_1} + r_k^{z_2} + r_k^v \quad (۶-۲۷)$$

$$r_k^\tau = -0.5(\tau(k) - \tau(k-1)) \quad \text{if } \tau > \tau_0 \quad (۶-۲۸)$$

¹⁰ منظور از 5^+ همه‌ی عددهای بزرگ‌تر یا مساوی ۵ است و منظور از 0^- همه‌ی عددهای کوچک‌تر یا مساوی ۰ است.

$$r_k^{z_1} = 0.1 \quad \text{if } \left[|z(k) - z_{\text{goal}}| < 0.2 \right] \wedge \left[(z(k) - z(k-1)) < 0.001 \right] \quad (6-29)$$

$$r_k^{z_2} = 5(z(k) - z(k-1)) \quad \text{if } z(k) < z_{\text{goal}} \quad (6-30)$$

$$r_k^v = -0.1 \quad \text{if } v(k) > v_{\text{max}} \quad (6-31)$$

برای یادگیری عامل از سیگنال تقویت (۶-۲۷) استفاده می‌کنم که خود از اجزای ساده‌تری تشکیل شده است. هر کدام از آن اجزا یکی از هدف‌های مساله را تامین می‌کند. به مانند پیش، سیگنال تقویت با استفاده از فرآیند سعی و خطا -و البته در نظر گرفتن دانش طراح نسبت به مساله- ایجاد شده است. سیگنال تقویت طراحی شده به سه جزء کارکردی قابل تقسیم است: جزءای که به کنترل زاویه‌ی کج‌شده‌گی می‌پردازد (۶-۲۸)، جزءای که مسوول کنترل ارتفاع جسم است ((۶-۲۹) و ((۶-۳۰) و جزءای که جلوی حرکت تند روبات‌ها را می‌گیرد (۶-۳۱). رابطه‌ی (۶-۲۸) با پاداش دادن (تنبیه کردن) کاهش (افزایش) زاویه‌ی کج‌شده‌گی به هنگامی که آن زاویه بزرگ‌تر از اندازه‌ی از پیش تعیین شده‌ی τ_0 باشد سعی در تقویت یادگیری کنترل زاویه می‌کند. رابطه‌ی (۶-۲۹) در ساکن ماندن نقطه‌ی تماس روبات-جسم را در نزدیکی هدف تشویق می‌کند. رابطه‌ی (۶-۳۰) بالا رفتن محرک روبات را به هنگامی که نقطه‌ی تماس روبات-جسم پایین‌تر از هدف است پاداش داده و پایین آمدن آن را تنبیه می‌کند. در نهایت، رابطه‌ی (۶-۳۱) حرکت تند روبات را تنبیه می‌کند. البته با توجه به ویژگی فضای عمل رفتارها، روبات‌ها هیچ‌گاه نمی‌توانند تا آن حدی سریع حرکت کنند که این مکانیزم فعال شود.

آخرین گام یادگیری هم‌زمان رفتار و ساختار است. سیگنال تقویت به کار رفته برای یادگیری ساختار و رفتار یک‌سان و برابر با سیگنال به کار رفته در یادگیری رفتار -با ساختار ثابت- است (رابطه‌ی (۶-۲۷)). در نتایج این بخش از روش یادگیری ساختار با بازنمایی مرتبه صفر استفاده شده است.

برای محک قرار دادن روش‌های یادگیری رفتار و یادگیری هم‌زمان رفتار/ساختار از دو معیار متوسط پاداش دریافت شده و معیار رفتارگرایانه‌ی (۶-۱۸) بهره می‌جوییم. مقایسه در دو فاز یادگیری و آزمون انجام می‌شود. در فاز اول، این دو معیار در اپیزودهای آغازین یادگیری که میزان کاوش نسبتاً زیاد است انجام می‌پذیرد و در فاز دوم این معیارها در شرایطی که چندین اپیزود از یادگیری گذشته و کاوش به پایان

رسیده بررسی می‌شوند. توجه کنید که در فاز آزمون یادگیری متوقف نمی‌شود و ساختار و یا رفتارها را ثابت نمی‌کنم – گرچه به دلیل وجود فاز یادگیری، در این فاز رفتارها و ساختار تقریباً یاد گرفته شده‌اند و تغییر قابل ملاحظه‌ای نمی‌کنند. با وجود آن که جزییات آزمایش در جدول ۱-۶ آمده است، اما اشاره به چند نکته مفید می‌نماید. ابتدا این که توزیع احتمال مکان‌های شروع حرکت در دو فاز یک‌سان است. در یادگیری هم‌زمان رفتار و ساختار، تغییر ساختار در هر سه اپیزود یک بار صورت می‌گیرد. این تغییر آرام‌تر ساختار به آن دلیل است تا اجازه دهیم رفتارها در طول آن سه اپیزود به ثبات برسند. به یاد داشته باشید که با تغییر ساختار، مقادیر ارزش حالت-عمل رفتارها ممکن است تغییر کند. در یادگیری ساختار صرف (بدون یادگیری رفتار) این تغییر در هر اپیزود صورت می‌گیرد.

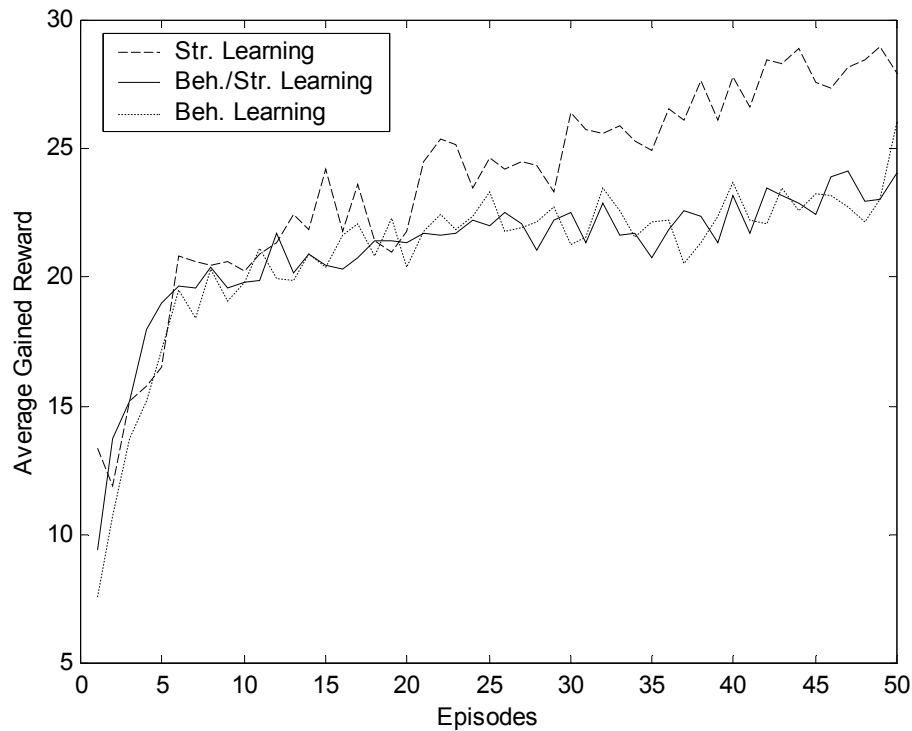
نکته‌ی مهم و قابل توجه دیگر این است که در طول آزمایش‌ها متوجه شدم که در بسیاری از موارد، عامل به خوبی وظیفه‌اش را یاد می‌گیرد اما در بعضی از موارد به پاسخ مناسب نمی‌رسد. دلایل ممکن برای این مشکل چندهدف بوده‌گی وظیفه، غیر-مارکوف بودن مساله از دید هر عامل به دلیل رویت‌پذیری جزئی و احتمالاً اشکال در طراحی سیگنال تقویت – که باعث گیر افتادن در پاسخ بهینه‌ی محلی می‌شود- می‌تواند باشد. در آن شرایط، عامل معمولاً پاسخ‌های ناقص‌ای چون یادگیری تنظیم زاویه بدون یادگیری توقف در کنار هدف، یادگیری توقف ولی حرکت نکردن به سمت بالا و ... را یاد می‌گرفت. با وجود آن که این پدیده مطلوب نیست، اما ممکن است بتوان مشکل وجود این عامل‌های درست یادنگرفته را با راه‌نمایی فعال عامل در طول یادگیری و بیش‌تر در معرض حالت‌های مشکل‌دار قرار دادن‌اش حل کرد^{۱۱}.

در شکل‌های در پیش رو مقایسه‌ای بین نتایج یادگیری رفتار، یادگیری هم‌زمان رفتار/ساختار و هم‌چنین یادگیری ساختار با رفتارهای از پیش طراحی‌شده (مشابه با آنچه در زیربخش پیش نمایش داده بودم)، معماری کاملاً توسط انسان طراحی شده (هم ساختار و هم رفتار) و معماری‌ای تصادفی (که رفتار و ساختار در آن به تصادف انتخاب شده‌اند) آورده می‌شود^{۱۲}.

^{۱۱} این ایده از کارهای ژان پیازه، روان‌شناس تکوینی، در یادگیری الهام گرفته شده است (به فصل مربوط به پیازه در [Hergenhahn01]

مراجعه کنید).

^{۱۲} ساختار از پیش مشخص طراحی شده مانند آن‌ای است که در آزمایش‌های طراحی ساختار آورده بودم و رفتارهای طراحی شده نیز مشابه آن‌ها هستند؛ یعنی:



شکل ۶-۲۰. (مساله‌ی بلندکردن جسم) مقایسه متوسط سیگنال تقویت دریافتی بین یادگیری ساختار (رفتار طراحی‌شده)، یادگیری رفتار (ساختار طراحی‌شده) و یادگیری همزمان ساختار و رفتار.

در شکل ۶-۲۰ منحنی یادگیری متوسط‌گیری‌شده بر روی صد اجرای مختلف چندین روش مختلف یادگیری در ۵۰ اپیزود اول نمایش داده شده است. دیده می‌شود که رفتار متوسط یادگیری رفتار و یادگیری همزمان رفتار/ساختار تا حدودی مشابه است اما منحنی متوسط یادگیری ساختار (با رفتارهای از پیش مشخص) بالاتر از آن دو قرار دارد. مقایسه‌ای دقیق بین این روش‌ها با استناد به اطلاعات آمارگان متوسط انجام‌پذیر نیست. به همین دلیل در شکل‌های بعدی از توزیع احتمال عمل کرد سیستم استفاده می‌شود. توزیع احتمال دانش بیش‌تری نسبت به فرآیند یادگیری در اختیارمان قرار می‌دهد.

Push more: $v(k+1) = 1$; **Do not go fast:** if $v(k) > v_{\max}$ then $v(k+1) = \max(0, v(k) - 0.5)$ else do nothing; **Stop at goal:** if $z(k) - z_{\text{goal}} > -0.2$ then stop ($v(k+1) = 0$); **Hurry up:** if $\tau(k) > \tau_0$ and the robot is the lowest one then $v(k+1) = 1$; **Slow down:** if $\tau(k) > \tau_0$ and the robot is the highest one then $v(k) = -1$.

توجه کنید که این رفتارها مشابه با رفتارهای (۶-۷)–(۶-۱۱) هستند اما به دلیل تفاوت فضای عمل‌شان تفاوت اندکی نیز بین این دو دسته رفتار وجود دارد.

شکل‌های ۶-۲۱ تا ۶-۲۴ توزیع احتمال تجمعی‌ی بیش‌تر بودن متوسط پاداش دریافتی و یا معیار عمل‌کرد رفتارگرایانه از میزان‌ای مشخص را نشان می‌دهند. برای نمونه، در شکل ۶-۲۱ احتمال این‌که عامل بیش از حد مشخص‌ای پاداش دریافت کند برای روش‌های مختلف یادگیری (و البته حالت تصادفی و از پیش طراحی‌شده) نمایش داده شده است. به عنوان نمونه‌ای برای چگونگی‌ی استفاده از این نمودار فرض کنید که می‌خواهیم درصد مواردی را که عامل‌مان با استفاده از یادگیری تقویتی به پاداش‌اش کم‌تر از ۲۰ می‌رسد پیدا کنیم. از نمودار مشخص است که حدود ۱۷٪ تجربه‌های یادگیری به پاداش‌ای کم‌تر از ۲۰ منجر می‌شود و ۸۳٪ نیز به پاداش‌اش بیش از ۲۰. احتمال دریافت متوسط پاداش کم‌تر از ۲۰ برای برای یادگیری هم‌زمان رفتار/ساختار ۳۰٪ است. مشخص است که هر چه دانش اولیه‌ی صحیح بیش‌تری داشته باشیم، احتمال رسیدن به پاسخ‌های به‌تر نیز بیش‌تر می‌شود. در نتیجه همان‌طور که در شکل ۶-۲۱ دیده می‌شود حالت‌ای که از پیش توسط انسان و با سعی و خطای بسیار طراحی شده است به دلیل آن‌که از ابتدا به بخش‌ای از دانش لازم برای حل مساله مجهز بوده است به پاسخ‌های بسیار به‌تری منجر شده است. هم‌چنین دیده می‌شود که معماری‌ای که در آن رفتارها و ساختار، هر دو، اتفاقی انتخاب شده‌اند بدتر از حالت‌های یادگرفته شده‌اند. عامل‌هایی که ساختارشان با روش یادگیری ساختار با بازنمایی مرتبه صفر طراحی شده به‌تر از عامل‌هایی هستند که رفتارهای‌شان یاد گرفته شده‌اند. نتایج یادگیری ساختار در مقادیر بزرگ پاداش به‌تر از یادگیری هم‌زمان است اما احتمال تولید عامل‌هایی با پاسخ‌های با کیفیت نه چندان مطلوب (کم‌تر از ۱۵) برای حالت هم‌زمان کم‌تر است. یادگیری رفتار صرف در مقادیر کم و زیاد پاداش بدتر از یادگیری هم‌زمان عمل می‌کند اما در مقادیر متوسط از یادگیری هم‌زمان به‌تر است. هم‌چنین احتمال وجود داشتن عامل‌هایی با پاداش بسیار زیاد، اندکی برای حالت هم‌زمان بیش از حالت یادگیری رفتار صرف است. این بدین دلیل است که یادگیری هم‌زمان توانسته است تعدادی ساختار تخصصی شده برای بعضی از مسایل پیدا کند که به‌تر از ساختار ثابت یادگیری رفتار صرف است. در شکل ۶-۲۲ همین نمودار برای فاز دوم، فاز آزمون، رسم شده است. در این فاز، عامل^۱ تجربه‌ی اپیزودهای زیادی را در چننه دارد و تقریباً نگاشت صحیح رفتارها و هم‌چنین ساختار مناسب‌اش را پیدا کرده است. دیده می‌شود که این نمودار نسبت به نمودار پیشین به سمت راست

تمایل پیدا کرده و تفاوت بین حالت‌های یادگیر و حالت طراحی شده کمتر شده است. مشابه همین نمودارها برای معیار عمل کرد رفتارگرایانه در شکل‌های ۶-۲۳ و ۶-۲۴ آورده شده است.

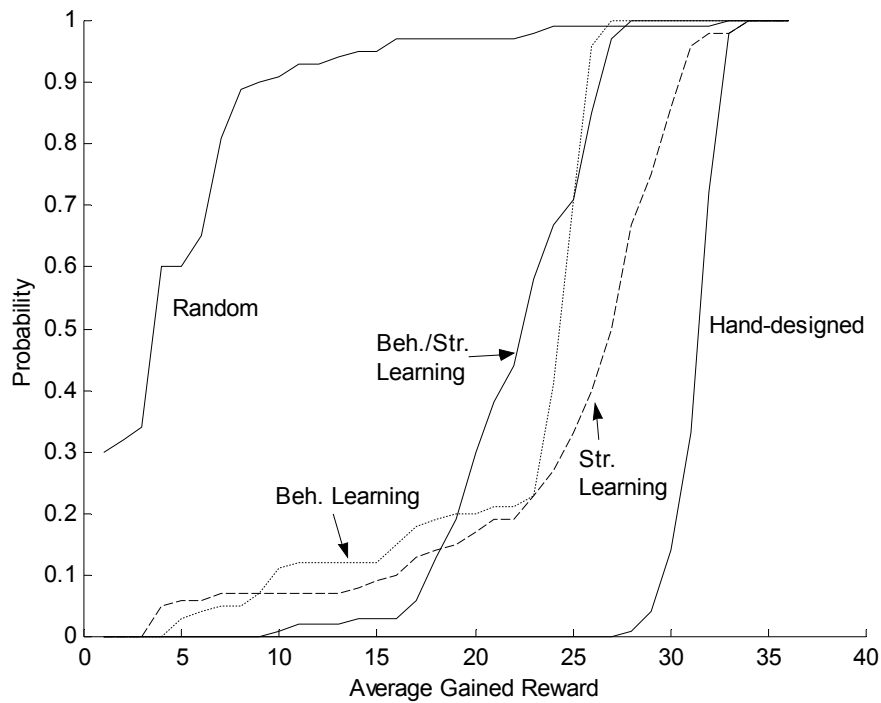
نمودارهای ۶-۲۵ تا ۶-۲۸ میانگین پاداش‌های دریافت شده (یا میانگین معیار رفتارگرایانه) برای درصد مشخص‌ای از عامل‌های برتر را نمایش می‌دهند.^{۱۳} مثلاً با انتخاب ۰,۳ بر روی محور افقی شکل ۶-۲۵ مشخص می‌شود که متوسط پاداش دریافتی ۷۰٪ برتر (حذف ۳۰٪ ضعیف‌تر) برابر با ۲۷ برای یادگیری ساختار و ۲۳ برای یادگیری هم‌زمان رفتار و ساختار است. اگر سمت چپ نمودارهای خیلی پایین نباشد مشخص می‌شود که تعداد حالت‌های خیلی بد محدود است. همچنین نموداری تقریباً مسطح نمایان‌گر پاسخ‌های در حدود هم خواهد بود: توزیع احتمال عمل کرد عامل‌های ایجاد شده نازک است. بدیهی است که معماری‌ای از پیش طراحی شده و ثابت به دلیل تغییر نکردن اجزای‌اش باعث ایجاد نموداری تقریباً مسطح می‌شود. شیب اندک نمودار ایجاد شده به دلیل نقاط آغازین تصادفی‌ی تماس روبات‌ها-جسم است که باعث تغییر اندکی در بازده می‌شوند. در شکل ۶-۲۵ دیده می‌شود که حالت از پیش طراحی شده از بقیه‌ی حالت‌ها بهتر است و یادگیری ساختار نیز برتر از یادگیری رفتار است. یک دلیل این پدیده کوچک‌تر بودن فضای فرضیه برای مساله‌ی یادگیری ساختار است - یادگیری ساختار مطلوب با فرض داشتن رفتارهای مناسب ساده‌تر از پیدا کردن رفتارهای مطلوب با فرض داشتن ساختار مناسب است. دلیل دیگر وجود مکانیزم انتخاب عمل کاوش‌گرایانه است که کیفیت عمل کرد عامل را در ابتدای یادگیری پایین می‌آورد. با صرف نظر کردن از انتخاب عمل کاوش‌گرایانه در فاز آزمون، تفاوت یادگیری ساختار و یادگیری رفتار کمتر می‌شود (شکل ۶-۲۶). در این فاز، با وجود آن که عمل کرد متوسط یادگیری رفتار پایین‌تر از عمل کرد متوسط یادگیری ساختار و معماری از پیش طراحی شده است، اما نتایج برترین عامل‌های ایجاد شده قابل مقایسه با نتایج آن دو است. به مانند قبل تعدادی (حدود ۴۰٪) از نتایج یادگیری هم‌زمان به‌تر از یادگیری رفتار صرف است و این به دلیل قابلیت انعطاف بیش‌تر ایجاد شده به دلیل یادگیری هم‌زمان رفتار و ساختار است. نکته‌ی قابل توجه دیگر تولید درصد قابل

^{۱۳} قابل ذکر است که نمودارهای ۶-۲۵ تا ۶-۲۸ از نمودارهای ۶-۲۱ تا ۶-۲۴ قابل استنتاج‌اند - چون نمودارهای دسته‌ی دوم بیان‌گر توزیع احتمال پدیده هستند و در نتیجه هر آمارگان دیگری نیز از آن‌ها قابل استخراج است - اما با این حال آوردن نمودارهای اخیر روشن‌گر خواهند بود.

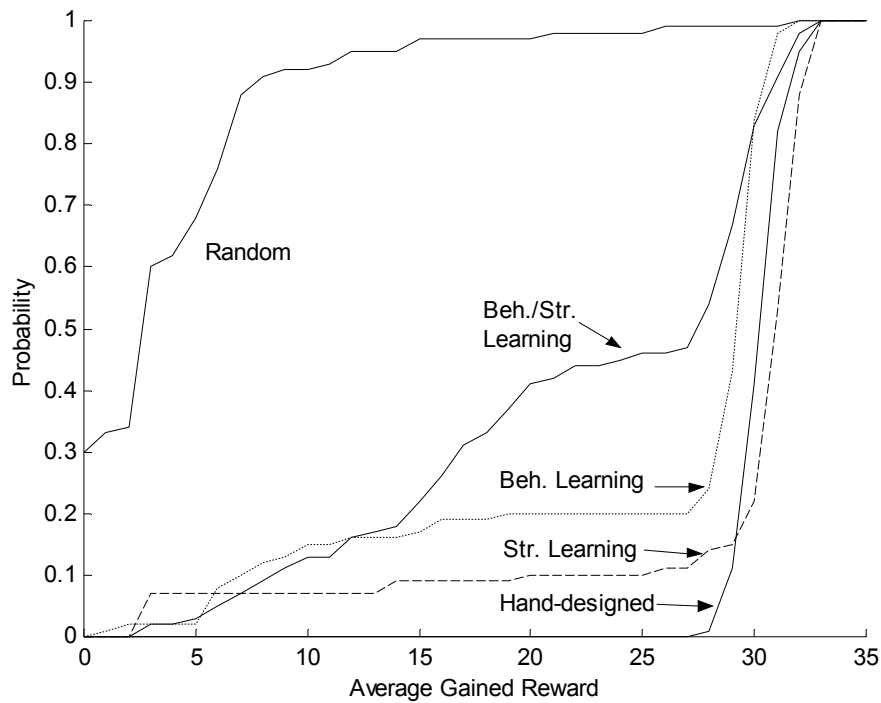
توجه‌ای (۹۰٪) عامل برتر از عامل‌های از پیش طراحی‌شده با استفاده از روش یادگیری ساختار است. نمودارهای ۶-۲۷ و ۶-۲۸ مشابه همین نوع نمودار برای معیار عمل کرد رفتارگرایانه است.

در نهایت، نمونه‌ای از مسیر نقاط اتصال جسم-روبات، زاویه‌ی کج‌شده‌گی جسم و رفتارهای کنترل‌کننده در طی عمل بلند کردن جسم را در شکل ۶-۲۹ آورده‌ام. این روبات‌ها عمل بلند کردن جسم را با استفاده از روش هم‌زمان یادگیری رفتار/ساختار آموخته‌اند. دیده می‌شود در ابتدای بلند کردن روبات‌های ۱ و ۲ ثابت‌اند و روبات ۳ رفتار "Hurry up" را انجام می‌دهد. سپس روبات ۱ شروع به انجام "Slow down" می‌کند و به سمت پایین حرکت می‌کند در حالی که روبات ۲ همچنان ثابت است و روبات ۳ نیز عمل پیشین، "Hurry up"، را انجام می‌دهد. پس از مدتی که زاویه‌ی کج‌شده‌گی به حد قابل قبول‌اش رسید، هر سه روبات شروع به هل دادن جسم به بالا (رفتار "Push More") می‌کنند تا به هدف برسند. پس از آن رفتار "Don't go fast" را انتخاب می‌کنند. این رفتار به گونه‌ای یاد گرفته شده است که در سرعت‌های پایین (سرعت‌ای که روبات عملاً در آن محدوده عمل می‌کند)، همچنان تصمیم به کم کردن سرعت نقاط اتصال روبات و جسم می‌کند. این کار با توجه به قید سرعت روبات‌ها (حداقل سرعت رفتار "Don't go fast" صفر است و جسم به کمک این رفتار نمی‌تواند به سمت پایین حرکت کند) باعث صفر شدن سرعت حرکت نقطه‌ی اتصال جسم-روبات می‌شود. در نتیجه به هنگامی که این رفتار کنترل‌کننده می‌شود، روبات توقف می‌کند.

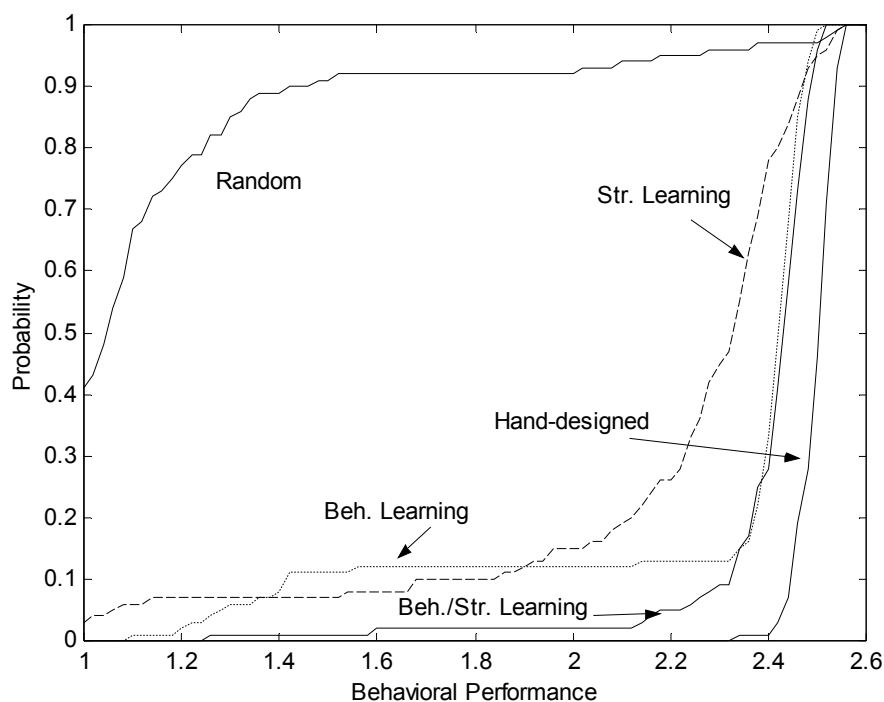
اشاره به این نکته ضروری است که نام‌گذاری یک رفتار (مثلاً رفتار "Don't go fast") الزاماً به معنای انجام رفتار مورد انتظار (مثلاً رفتار سریع حرکت نکردن) نیست. در فرآیند یادگیری، رفتارها به گونه‌ای یاد گرفته می‌شوند که منجر به افزایش بازده عامل -از دیدی عینی و مطابق با تابع تقویت- می‌شود که این در همه‌ی موارد معادل با انتظار ذهنی‌ی طراح (مثلاً کم شدن سرعت‌های زیاد) نیست. در نهایت می‌توان این‌گونه جمع‌بندی کرد که روش‌های یادگیری معرفی شده برای طراحی خودکار عامل‌های رفتارگرا می‌توانند به طراحی‌هایی قابل رقابت -و حتی به‌تر- از طراحی‌های انسانی بیانجامند.



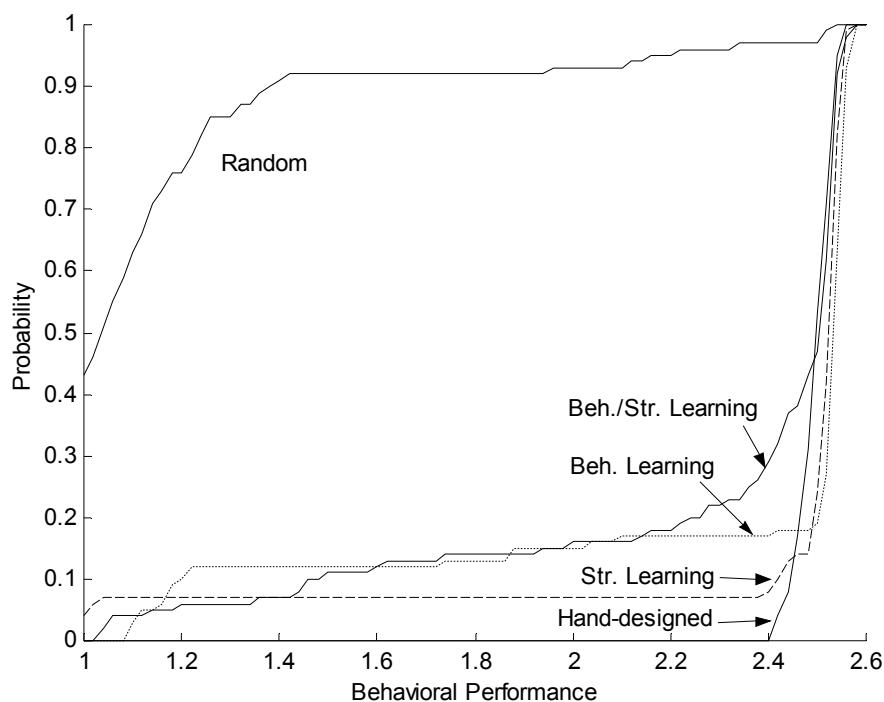
شکل ۶-۲۱. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی سیگنال تقویت دریافتی ($P\{\text{average gained reward} \leq x\}$) در فاز یادگیری (بین ایزود ۱ تا ۵۰. کاوش برقرار است) به ازای ترکیب‌های مختلف روش‌های یادگیری.



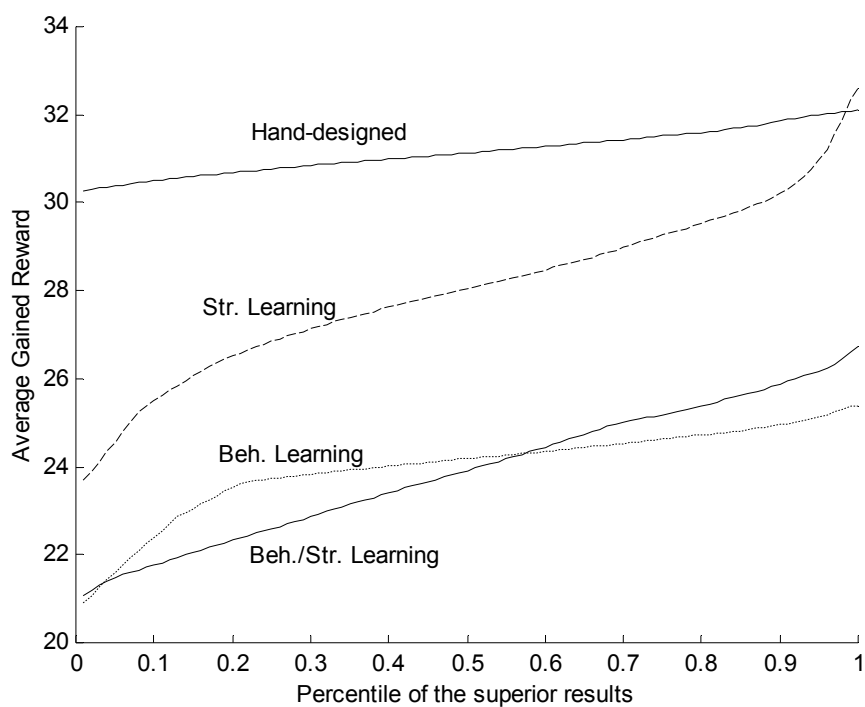
شکل ۶-۲۲. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی سیگنال تقویت دریافتی ($P\{\text{average gained reward} \leq x\}$) در فاز آزمون (بین ایزود ۵۱ تا ۱۰۰. کاوش پایان یافته است) به ازای ترکیب‌های مختلف روش‌های یادگیری.



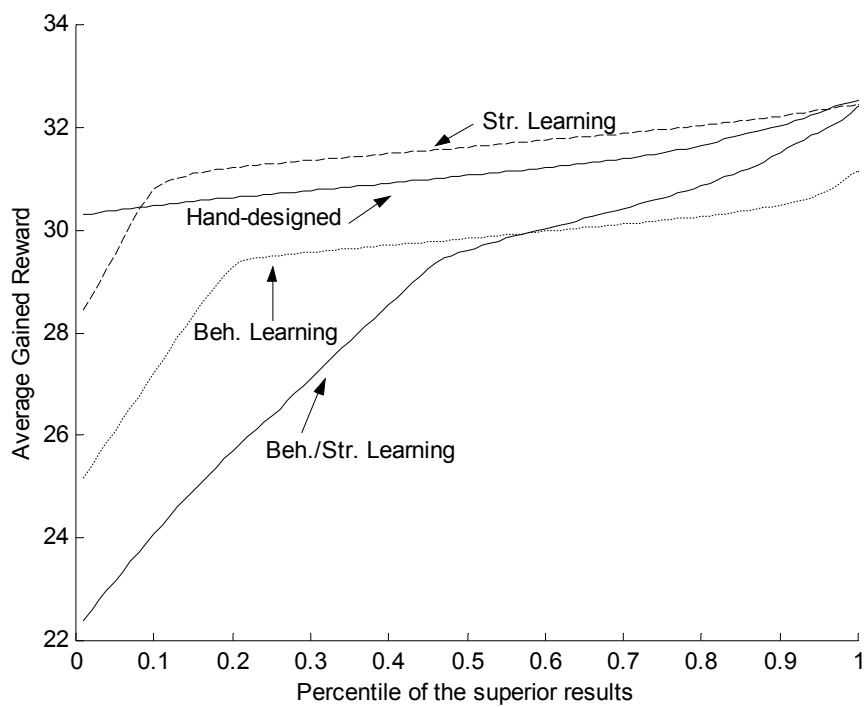
شکل ۲۳-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی معیار رفتارگرایانه ($P\{\text{behavioral performance} \leq x\}$) در فاز یادگیری (بین ایزود ۱ تا ۵۰ کاوش برقرار است) به ازای ترکیب‌های مختلف روش‌های یادگیری.



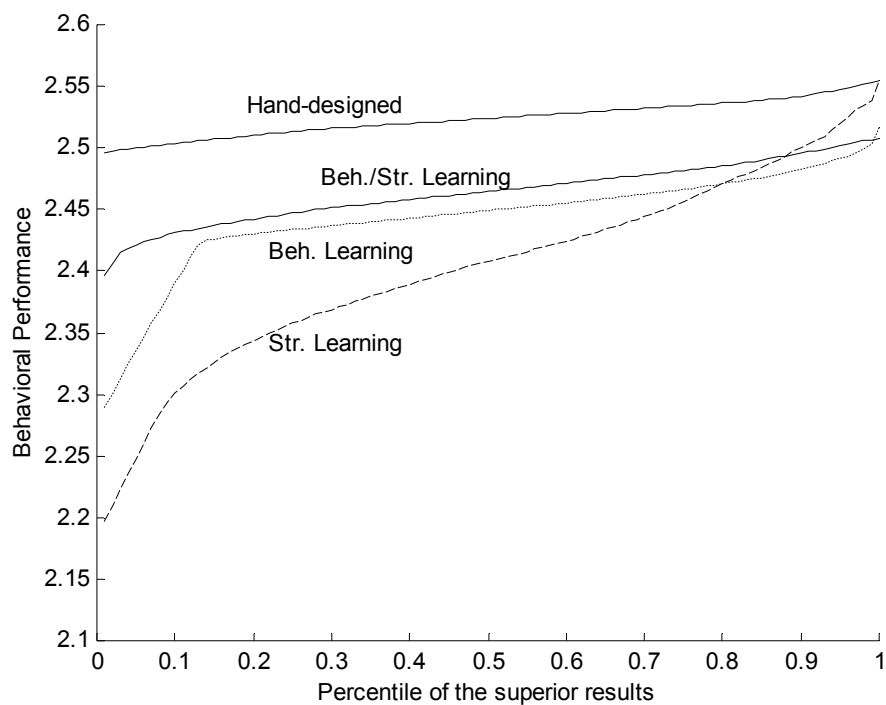
شکل ۲۴-۶. (مساله‌ی بلندکردن جسم) مقایسه توزیع تجمعی معیار رفتارگرایانه ($P\{\text{behavioral performance} \leq x\}$) در فاز آزمون (بین ایزود ۵۱ تا ۱۰۰ کاوش پایان یافته است) به ازای ترکیب‌های مختلف روش‌های یادگیری.



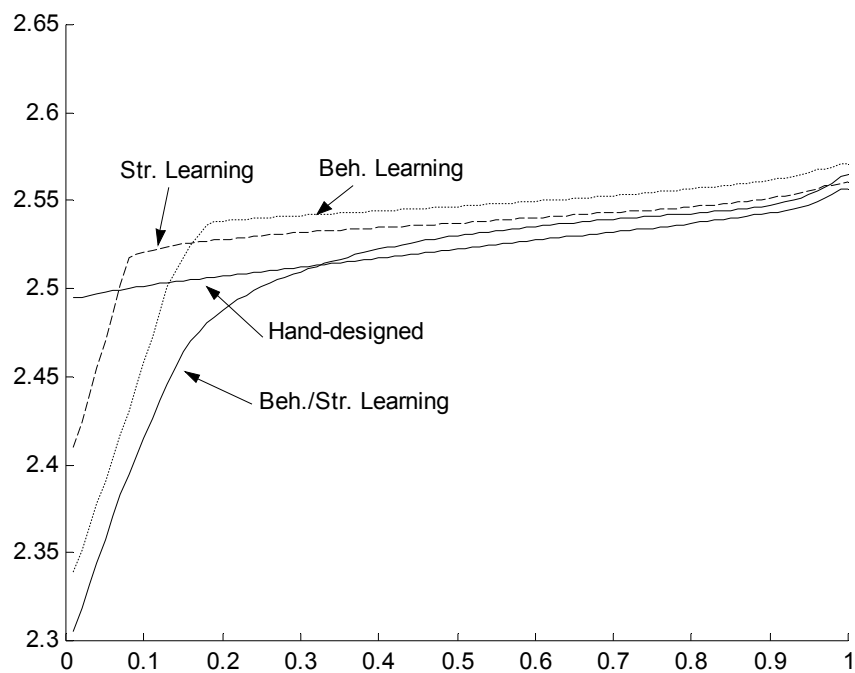
شکل ۲۵-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی پاداش دریافت شده در فاز یادگیری (اپیزود ۱ تا ۵۰. کاوش برقرار است) به ازای روش‌های مختلف یادگیری.



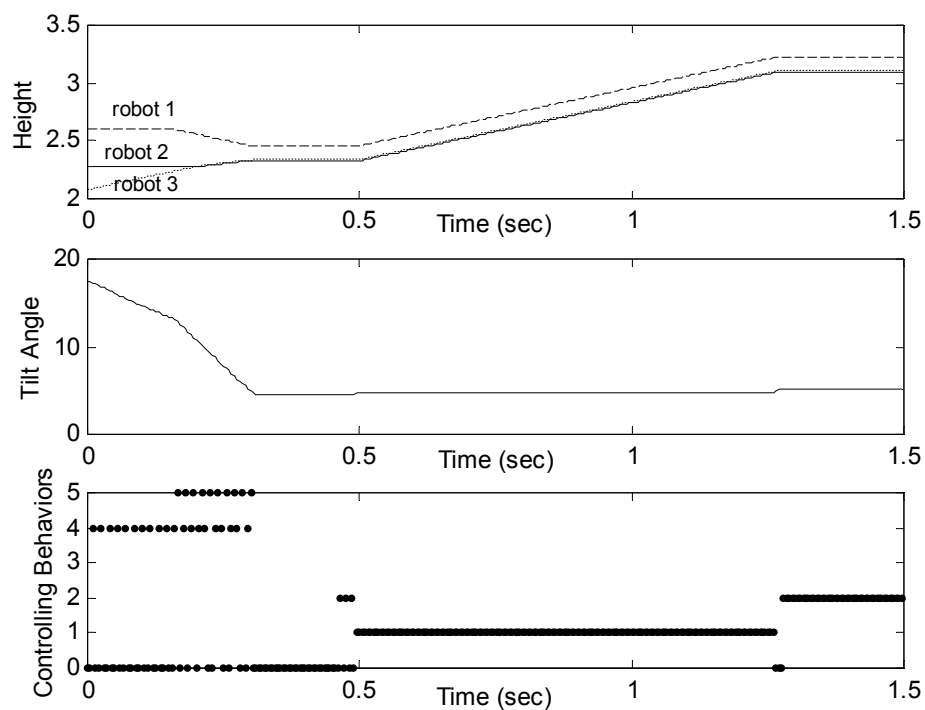
شکل ۲۶-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی پاداش دریافت شده در فاز آزمون (اپیزود ۵۱ تا ۱۰۰. کاوش پایان یافته است) به ازای روش‌های مختلف یادگیری.



شکل ۲۷-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی معیار رفتارگرایانه در فاز یادگیری (اپیزود ۱ تا ۵۰. کاوش برقرار است) به ازای روش‌های مختلف یادگیری.



شکل ۲۸-۶. (مساله‌ی بلندکردن جسم) متوسط تجمعی معیار رفتارگرایانه در فاز آزمون (اپیزود ۵۱ تا ۱۰۰. کاوش پایان یافته است) به ازای روش‌های مختلف یادگیری.



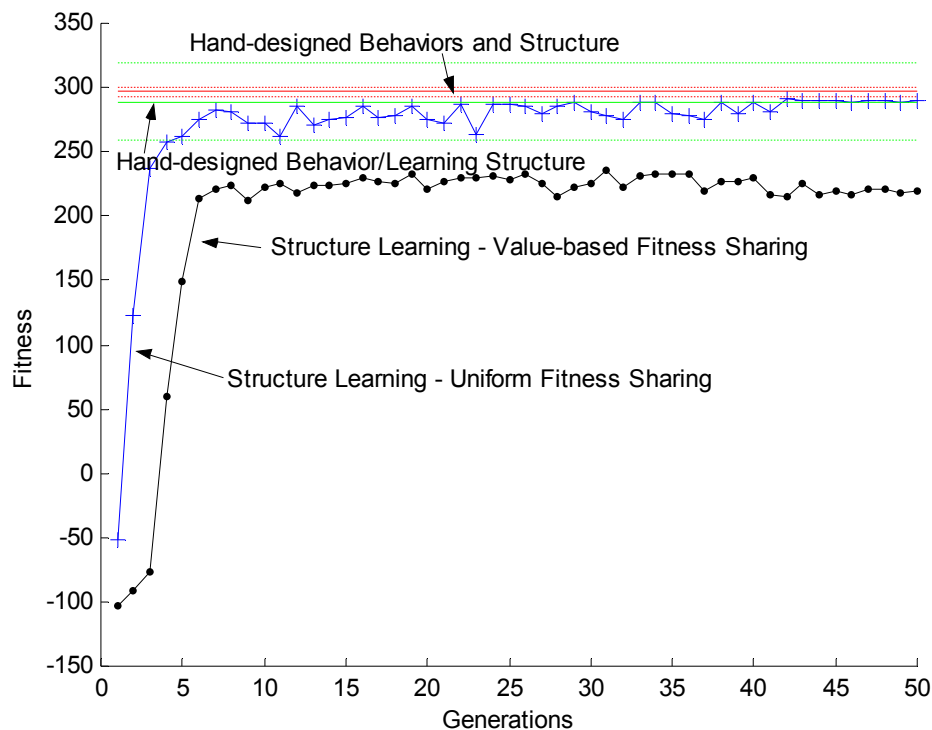
شکل ۲۹-۶. (مساله‌ی بلندکردن جسم) نمونه‌ای از مسیر حرکت نقاط تماس روبات-جسم، زاویه‌ی کج‌شده‌گی به هنگام بلند کردن جسم، و رفتارهای کنترل‌کننده‌ی روبات در هر واحد زمانی. این مسیر حاصل چندین اپیزود یادگیری هم‌زمان رفتار و ساختار است. تناظر رفتارها و شماره‌های روی نمودار پایین بدین گونه است:

0 (No Behavior), 1 (Push More), 2 (Don't Go Fast), 3 (Stop), 4 (Hurry up), 5 (Slow down).

۶-۲-۳) تکامل رفتار و یادگیری ساختار

در این بخش از ایده‌های تکاملی مطرح شده در فصل ۵ برای حل مسأله‌ی بلند کردن جسم توسط چند روبات بهره می‌جویم. به مانند پیش، تکامل رفتار را در دو حالت اشتراک سازگاری و یک‌نواخت و ارزش‌محور بررسی می‌کنم. علاوه بر این، تاثیر استفاده از ایده‌های الگوریتم ممیتیک - که در اعضای جدید جامعه از ابتدای کنش‌شان با محیط به الگوهای فرهنگی پیشینیان‌شان مجهز هستند - نیز مطالعه می‌شود. در این مسأله، دو معیار بررسی می‌شود: (۱) سازگاری و چند اپیزود نهایی عامل که معادل با متوسط سیگنال تقویت دریافتی در چند اپیزود نهایی است و (۲) سازگاری کل دوره‌ی زندگی عامل که متوسط سیگنال تقویت دریافتی از ابتدا تا انتهای زندگی عامل است. معیار اول مشخص می‌کند که عامل در نهایت به چه وضعیت‌ای رسیده است؛ در حالی که معیار دوم بیان‌گر کل کیفیت کار عامل در طول زندگی‌اش است. مطلوب است که این دو معیار سازگاری تا حد ممکن بزرگ و به هم نزدیک باشند. در این صورت عامل در همه‌ی طول عمرش رفتار قابل قبول‌ای داشته است. اگر معیار دوم بسیار کوچک‌تر از معیار اول باشد، عامل در ابتدای زندگی تلاش‌های ناموفق‌ای برای یادگیری ساختار مناسب داشته است ولی به تدریج روش موفق حل مسأله را فراگرفته است.

سازگاری عامل‌ها با توجه به نوع اشتراک سازگاری مشخص می‌شوند. در اشتراک یک‌نواخت، سازگاری عامل در چند اپیزود انتهایی به طور یک‌سان به همه‌ی رفتارها داده می‌شود. در اشتراک ارزش‌محور، $V_{T(B_t)}$ به هر عامل داده می‌شود که برای محاسبه‌اش از همان پارامتر $\alpha_{episode}^{str}$ ای که در یادگیری ساختار عامل به کار رفته استفاده شده است. با توجه به نرخ کاهش $\alpha_{episode}^{str}$ ، این تابع ارزش بیش‌تری بر آخرین اپیزودها می‌گذارد. در هر حال برای نمایش عمل‌کرد سیستم، از دو معیار بیان شده استفاده می‌شود. پس باید به یاد داشته باشیم که معیار بهینه‌سازی و آنچه نمایش داده می‌شود دقیقاً یک‌سان نیست. این تفاوت برای حالت اشتراک یک‌نواخت تنها در معیار متوسط سازگاری در کل زندگی عامل است در حالی که برای اشتراک ارزش‌محور، هر دو معیار نمایش داده شده با معیار بهینه‌گی تفاوت دارند.



شکل ۶-۳۰. (مساله‌ی بلند کردن جسم) متوسط سازگاری پنج اپیزود نهایی به ازای طراحی‌های مختلف: (۱) تکامل رفتارها (اشتراک سازگاری یک‌نواخت) و یادگیری ساختار (آبی، ۲) تکامل رفتارها (اشتراک سازگاری ارزش‌محور) و یادگیری ساختار (سیاه، ۳) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۴) رفتارها و ساختار طراحی‌شده (قرمز). خطوط نقطه‌چین اطراف حالت‌های طراحی‌شده (۳ و ۴) ناحیه‌ی یک انحراف معیار از متوسط را نشان می‌دهند.

در شکل ۶-۳۰ مقایسه‌ای بین سازگاری متوسط جمعیت در چند حالت مختلف نشان داده شده است. در آن شکل، نتیجه‌ی حاصل از عامل‌هایی که رفتارشان توسط یکی از دو روش تکاملی (اشتراک سازگاری یک‌نواخت و ارزش‌محور) به وجود آمده‌اند و ساختار را نیز یاد می‌گیرند، آن‌هایی که رفتارهای از پیش طراحی‌شده و با یادگیری ساختار دارند (با رفتارهای از پیش طراحی‌شده‌ای مطابق با آنچه در قسمت یادگیری ساختار معرفی کردم) و همچنین عامل‌هایی که هم رفتار و هم ساختارشان (درست همان ساختار معرفی شده در بخش یادگیری ساختار) از پیش طراحی‌شده است با هم مقایسه شده است. مقایسه بر روی سازگاری حاصل از متوسط سیگنال تقویت پنج اپیزود نهایی یادگیری بوده است و در حالت‌ای که یادگیری‌ای وجود نداشته، هر آزمون از عامل تنها پنج اپیزود اجرا می‌شده است. برای جزییات بیش‌تر به جدول ۶-۲ مراجعه کنید.

در شکل ۶-۳۰ مشاهده می‌شود که متوسط سازگاری عامل یادگیر که رفتارهای اش با استفاده از اشتراک سازگاری یک‌نواخت تکامل می‌یابند بسیار سریع به حد طراحی رفتار و ساختار از پیش تعیین شده و طراحی رفتار از پیش تعیین شده همراه با یادگیری ساختار می‌رسد. حالت دیگر، عامل یادگیر ساختار همراه با تکامل رفتار از نوع اشتراک سازگاری ارزش‌محور، دارای متوسط پایین‌تری است. قابل ذکر است که هر دوی این روش‌ها عامل‌هایی تولید می‌کنند که سازگاری‌ای در حد به‌ترین نتایج روش از پیش طراحی شده دارند ولی برای جلوگیری از شلوغی شکل و به این دلیل که آن عامل‌ها از همان نسل‌های اولیه به وجود می‌آیند و در طول تکامل وجودشان (و نه تعدادشان) تغییری نمی‌کند نشان داده نشده‌اند. برخلاف مساله‌ی مجرد، این مساله دارای رفتارهایی نیست که بهینه‌گی هر کدام معادل با بهینه‌گی کل سیستم باشد. مثلاً اگر نتیجه‌ی تکامل و یادگیری این باشد که معماری‌ای با ساختار $T = [... \text{PushMore} \text{ Stop}]$ با رفتارهای “Push more” و “Stop” مطابق آن‌چه انتظار می‌رود (یعنی ۶-۲۲ و ۶-۲۴) ایجاد شود، این دو رفتار هر دو پاداش بسیار زیادی می‌گیرند؛ ولی این ترکیب به‌ترین ممکن برای عامل نیست چون از رفتارهای تنظیم زاویه هیچ استفاده‌ای نشده است. در نتیجه، عامل‌های خودخواه‌ای^{۱۴} که حاصل این نوع اشتراک سازگاری هستند الزاماً به به‌ترین پاسخ عامل نمی‌انجامند. اشتراک سازگاری یک‌نواخت با این که در اشتراک سازگاری چندان هوش‌مندانه عمل نمی‌کند اما به نظر می‌رسد که برای این نوع مسایل به دلیل طبیعت تیم‌نگر خود به‌تر عمل می‌کند.

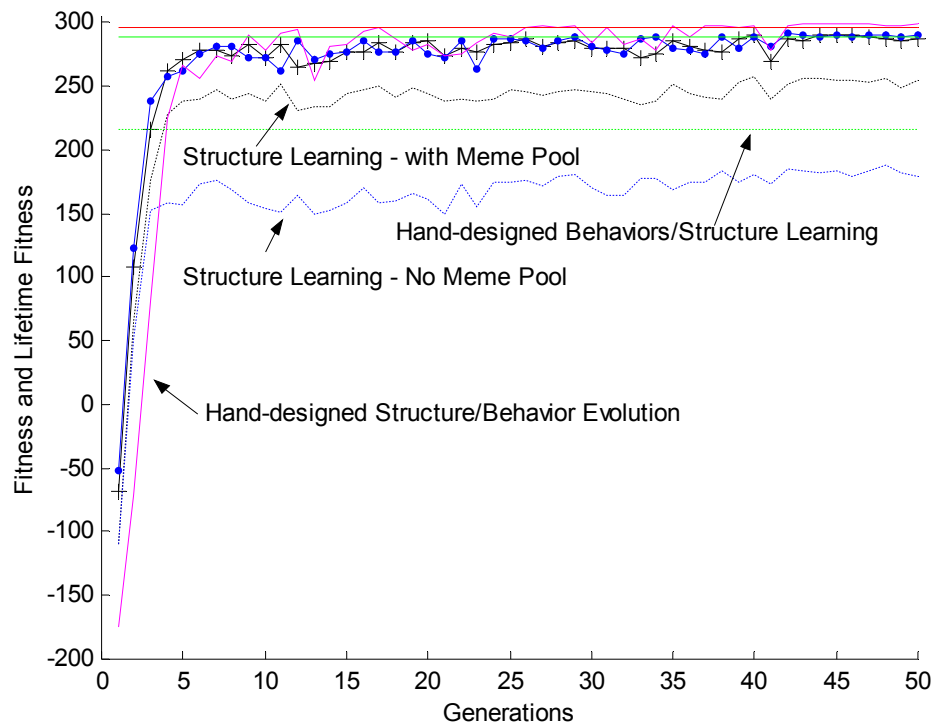
اینک بررسی‌ی دقیق‌تری روی هر کدام از این روش‌های تکامل و تاثیر یادگیری و الگوریتم ممیتیک انجام می‌دهم. نتایج آزمون روش اشتراک سازگاری یک‌نواخت در حالت‌های مختلف در شکل ۶-۳۱ آمده است. مشاهده می‌شود که تفاوت محسوس‌ای بین متوسط سازگاری پنج اپیزود نهایی (خطوط پررنگ) حالت‌های ساختار از پیش مشخص، یادگیری ساختار بدون کمک مم و یادگیری ساختار با اطلاع اولیه‌ی منتقل شده توسط مم وجود ندارد. همه‌ی حالت‌ها منجر به نتایج قابل رقابت با ساختار/رفتار طراحی شده و یا رفتارهای طراحی شده/ساختار یادگرفته شده می‌شوند. اما مقایسه‌ی بین متوسط سازگاری در طول دوره‌ی زندگی‌ی عامل (خطوط نقطه‌چین) نتایج جالبی را نمایان می‌کند. وقتی

¹⁴ Selfish

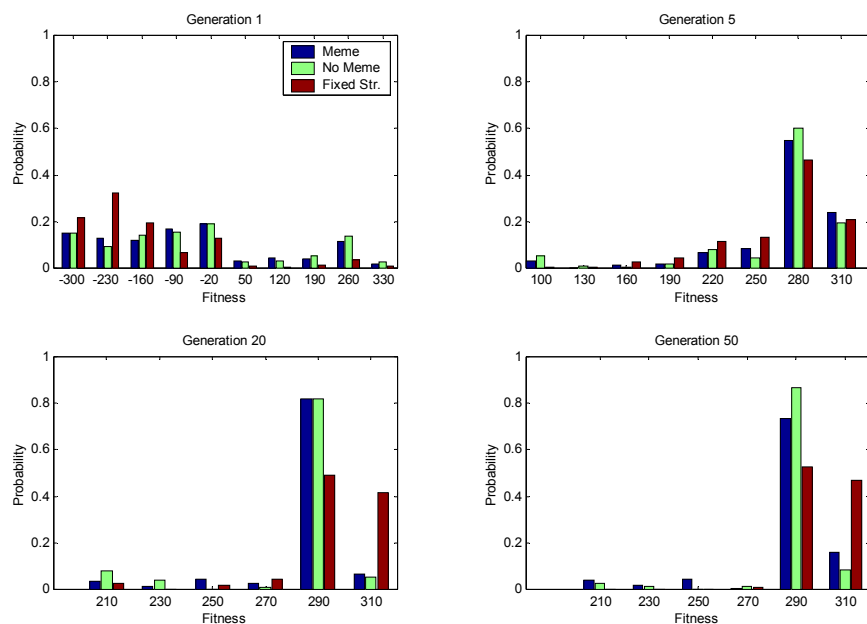
از الگوریتم مم‌تیک استفاده می‌کنیم، میزان سازگاری کل زندگی بسیار بالاتر از هنگامی است که از مم استفاده نکرده‌ایم. حتی این میزان از وقتی که طراحی رفتارها از قبل مشخص است و تنها ساختارشان یادگرفته می‌شود نیز بیش‌تر است و فاصله‌ی بین سازگاری نهایی عامل و سازگاری کل زندگی به میزان قابل توجه‌ای کم‌تر از حالت بدون مم است. در شکل ۳۲-۶، هیستوگرام سازگاری پنج اپیزود نهایی در چند نسل مختلف نمایش داده می‌شود. مشخص است که تفاوت آن‌چنانی‌ای بین حالت با مم و حالت بدون مم نیست (جز این‌که به نظر می‌رسد حالت با مم می‌تواند اندکی عامل بسیار خوب بیش از حالت بدون مم تولید کند)، اما در شکل ۳۳-۶ که هیستوگرام سازگاری کل دوره زندگی عامل است تفاوت این دو حالت بسیار واضح است: الگوریتم مم‌تیک شانس تولید عامل‌های خوب را بسیار بالا می‌برد. نمودار توزیع تجمعی سازگاری عامل‌های ایجاد شده (برای هر دو تعریف سازگاری) در شکل ۳۴-۶ آمده است. دیده می‌شود که تفاوت بین سازگاری نهایی (خطوط پررنگ) و سازگاری کل زندگی (خطوط نقطه‌چین) وقتی از ایده‌ی انتقال مم بهره می‌گیریم، تا حد زیادی کم می‌شود.

نتایج آزمون روش اشتراک سازگاری ارزش‌محور در حالت‌های مختلف در شکل ۳۵-۶ آمده است. متوسط سازگاری نهایی جمعیت به ازای همه‌ی روش‌ها -یعنی تکامل با ساختار ثابت، یادگیری بدون مم و یادگیری با مم- تقریباً یک‌سان است و تنها سرعت هم‌گرایی به پاسخ مناسب برای حالت ساختار ثابت کمی بیش‌تر است. مطابق قبل تفاوت سازگاری کل زندگی برای وقتی که از الگوریتم مم‌تیک استفاده می‌کنیم کم‌تر از زمان‌هایی است که از آن بهره نمی‌جوییم. گرچه تفاوت بین دو معیار سازگاری به نسبت کم‌تر از وقتی است که از اشتراک سازگاری یک‌نواخت استفاده شده بود. در شکل ۳۶-۶ مشاهده می‌شود که سازگاری نهایی حالت‌ای که از الگوریتم مم‌تیک استفاده شده است در نهایت (نسل ۲۰۰) منجر به درصد بالاتری از عامل‌های برتر می‌شود و این تفاوت وقتی سازگاری کل زندگی عامل معیار قضاوت‌مان باشد، بسیار واضح‌تر است (شکل ۳۷-۶). توزیع تجمعی سازگاری را در شکل ۳۸-۶ آورده‌ام. برتری سازگاری کل زندگی برای حالت استفاده از مم نسبت به حالت بدون مم واضح است. نکته‌ی جالب -و متفاوت با اشتراک سازگاری یک‌نواخت- در این است که این تفاوت نه تنها در سازگاری کل زندگی که در سازگاری نهایی نیز دیده می‌شود. هم‌چنین تفاوت کم‌تری بین دو سازگاری نسبت به روش قبلی اشتراک سازگاری وجود دارد.

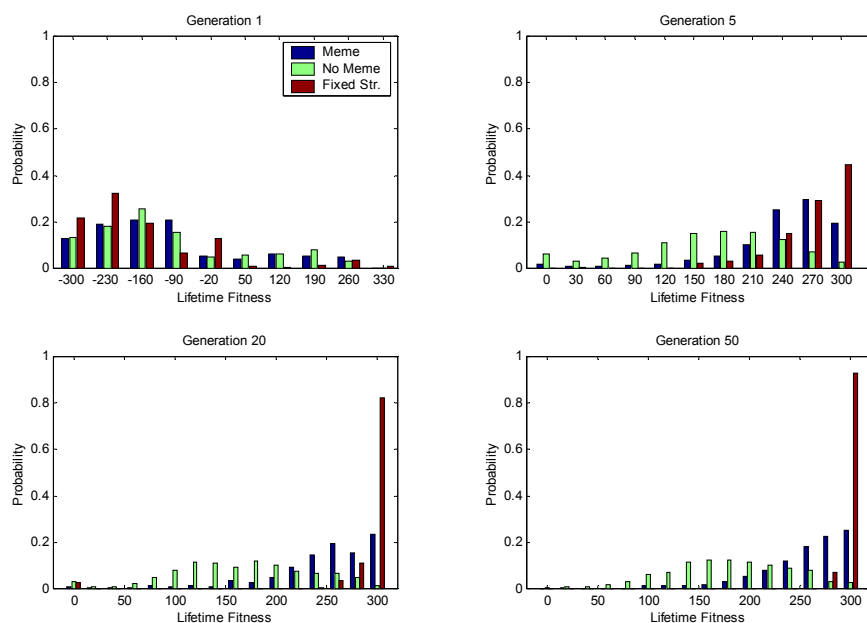
در نهایت می‌توان این‌گونه جمع‌بندی کرد که روش‌های تکامل معرفی شده برای طراحی خودکار عامل‌های رفتارگرا می‌توانند به طراحی‌هایی قابل رقابت – و حتی بهتر – از طراحی‌های انسانی بیانجامند. هم‌چنین استفاده از ایده‌ی الگوریتم ممیتیک می‌تواند باعث افزایش کیفیت عامل از ابتدای زندگی‌اش شده و یادگیری را بسیار ساده‌تر کند. در این مساله، اشتراک سازگاری ی یک‌نواخت بهتر از اشتراک سازگاری ارزش‌محور عمل کرد.



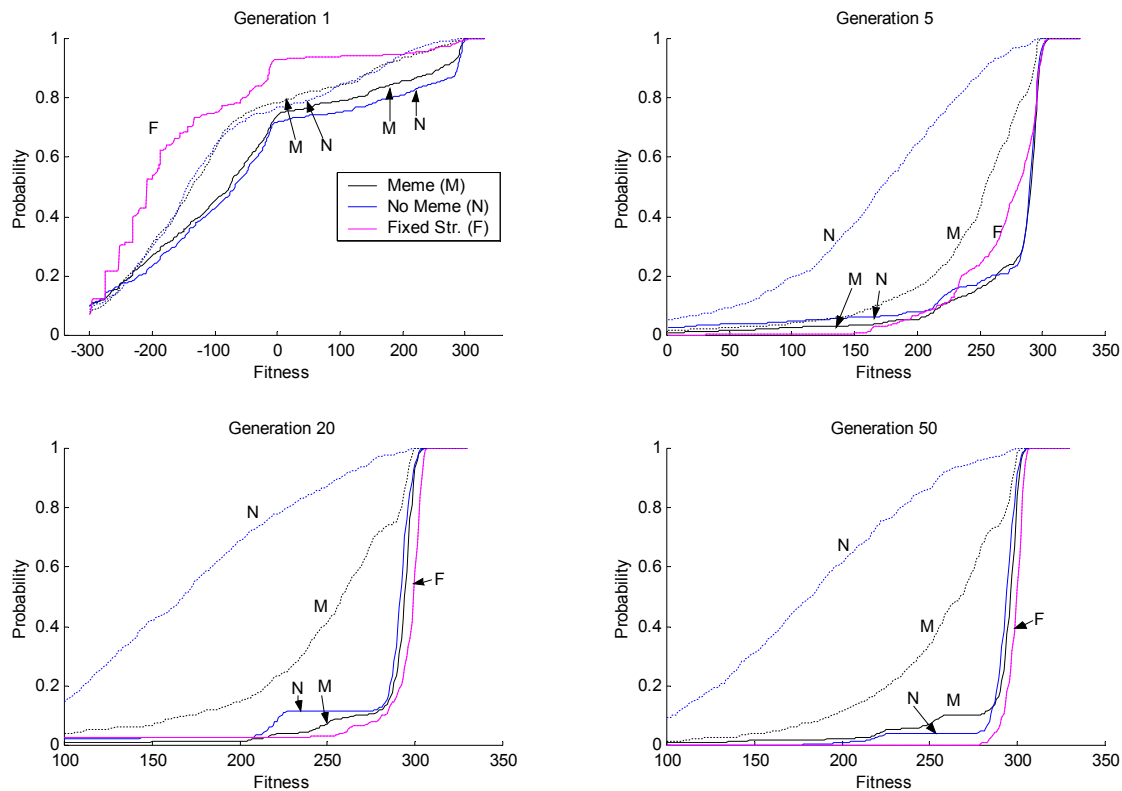
شکل ۳۱-۶. (مساله‌ی بلند کردن جسم) متوسط سازگاری پنج ایزود نهایی و متوسط سازگاری کل دوره‌ی زندگی عامل در طراحی‌های مختلفی که از اشتراک سازگاری یک‌نواخت استفاده می‌کنند: (۱) تکامل رفتارها و یادگیری ساختار بدون کمک مم (آبی)، (۲) تکامل رفتارها و یادگیری ساختار با کمک مم (سیاه)، (۳) تکامل رفتارها و ساختار طراحی‌شده (بنفش)، (۴) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۵) رفتارها و ساختار طراحی‌شده (قرمز). خطوط پررنگ بیان‌گر متوسط سازگاری در پنج ایزود انتهایی زندگی عامل هستند و خطوط نقطه‌چین سازگاری کل زندگی عامل را نشان می‌دهند. با وجود آن‌که عمل‌کرد آخرین ایزودهای همگی عامل‌ها تقریباً یک‌سان است اما عمل‌کرد کل دوره‌ی زندگی عامل‌هایی که از اطلاعات مم استفاده می‌کنند بسیار بالاتر است.



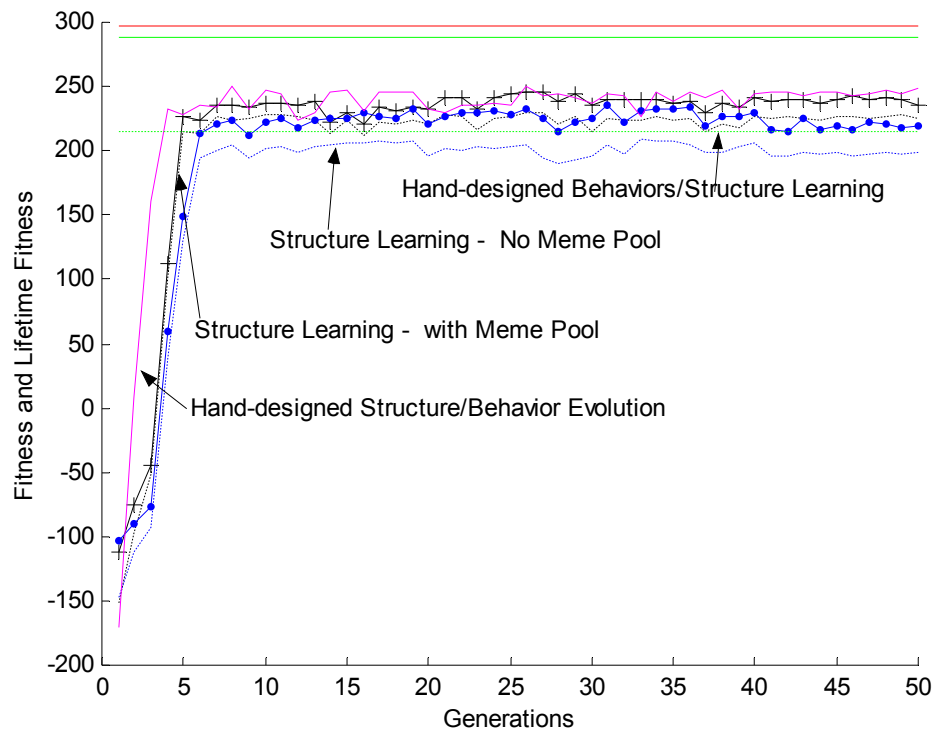
شکل ۳۲-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.



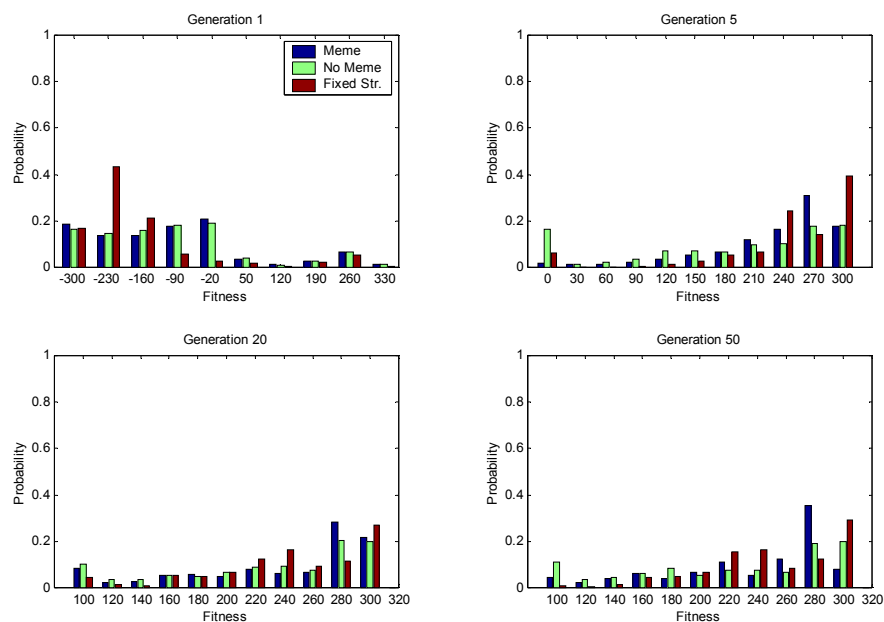
شکل ۳۳-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری کل دوره‌ی زندگی عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.



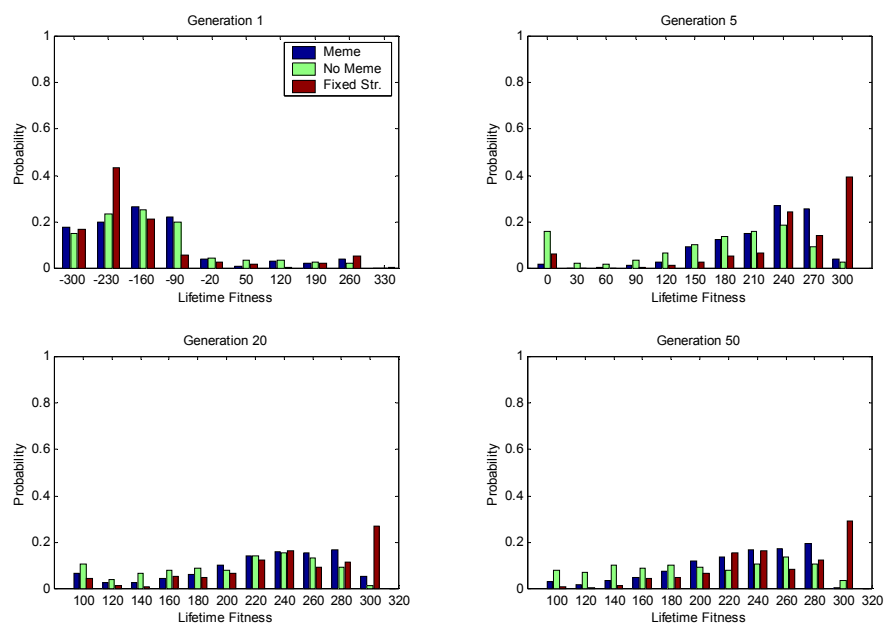
شکل ۳۴-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی احتمال تجمعی سازگاری ($P\{Fitness \leq f\}$) عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری یک‌نواخت. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (سیاه)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (آبی) و عامل‌هایی که از ساختار طراحی شده استفاده کرده‌اند (بنفش) انجام شده است. خطوط نقطه‌چین سازگاری کل دوره‌ی زندگی عامل را نشان می‌دهند. هر چه توزیع بیش‌تر به سمت راست گرایش داشته باشد، شانس ایجاد عامل‌های خیلی خوب بیش‌تر است.



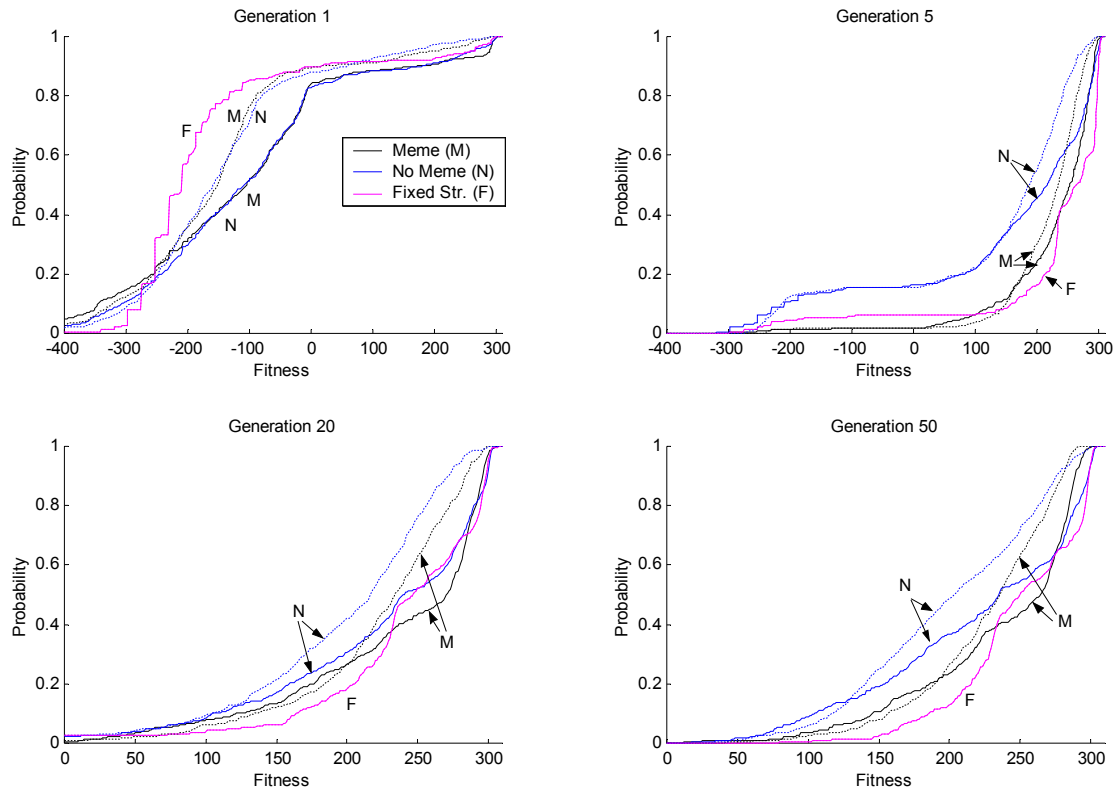
شکل ۳۵-۶. (مساله‌ی بلند کردن جسم) متوسط سازگاری پنج ایزود نهایی و متوسط سازگاری کل دوره‌ی زندگی عامل در طراحی‌های مختلفی که از اشتراک سازگاری ارزش‌محور استفاده می‌کنند: (۱) تکامل رفتارها و یادگیری ساختار بدون کمک مم (آبی)، (۲) تکامل رفتارها و یادگیری ساختار با کمک مم (سیاه)، (۳) تکامل رفتارها و ساختار طراحی‌شده (بنفش)، (۴) رفتارهای طراحی‌شده و یادگیری ساختار (سبز) و (۵) رفتارها و ساختار طراحی‌شده (قرمز). خطوط پررنگ بیان‌گر متوسط سازگاری در پنج ایزود انتهایی زندگی عامل هستند و خطوط نقطه‌چین سازگاری کل زندگی عامل را نشان می‌دهند. با وجود آن‌که عمل‌کرد آخرین ایزودهای همگی عامل‌ها تقریباً یک‌سان است اما عمل‌کرد کل دوره‌ی زندگی عامل‌هایی که از اطلاعات مم استفاده می‌کنند بالاتر است.



شکل ۳۶-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.



شکل ۳۷-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی چگالی احتمال سازگاری کل دوره‌ی زندگی عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (آبی تیره)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (سبز) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (قرمز) انجام شده است.



شکل ۳۸-۶. (مساله‌ی بلندکردن جسم) مقایسه‌ی احتمال تجمعی سازگاری ($P\{Fitness \leq f\}$) عامل‌های ایجاد شده با روش تکاملی با اشتراک سازگاری ارزش‌محور. مقایسه بین حالت‌هایی که عامل‌ها از اطلاعات مم استفاده کرده‌اند (سیاه)، عامل‌هایی که از ساختاری تصادفی شروع کرده‌اند (آبی) و عامل‌هایی که از ساختار طراحی‌شده استفاده کرده‌اند (بنفش) انجام شده است. خطوط نقطه‌چین سازگاری کل دوره‌ی زندگی عامل را نشان می‌دهند. هر چه توزیع بیش‌تر به سمت راست گرایش داشته باشد، شانس ایجاد عامل‌های خیلی خوب بیش‌تر است.

	Abstract Problem (Structure Learning)	Abstract Problem (Behavior Learning)	Object Lifting (Structure Learning)	Object Lifting (Behavior(/Structure) Learning)
Problem Space	50x50 problem space and two actions	50x50 problem space and four actions	Multi-robot object lifting	Multi-robot object lifting
Solution Space	Four behaviors $B_i (i = 1,2,3,4)$ (see Fig. 3)	Four behaviors each of $B_i (i = 1,2,3,4)$ and a_i them consists of NA	Five behaviors, i.e. <i>Push more, Don't go fast, Stop at goal, Hurry up, and Slow down.</i>	Five naïve behaviors' state and action space without any specified mapping.
Learning Rates	$\alpha_{episode} =$ $\alpha_0 (0.999)^{episode}$ $\alpha_0 = 0.05$ (ZO) $\alpha_0 = 0.1$ (FO)	$\alpha_{episode}^{beh} = \alpha_{episode}^{str} =$ $\frac{0.2}{1 + \text{mod}(episode, 1000)}$ ϵ -greedy action selection with $\epsilon_{episode} = \frac{2}{\ln(episode + 1)}$	$\alpha_{episode}^{str} = 0.1(0.99)^{episode}$	$\alpha_{episode}^{beh} = 0.1$ $\alpha_{episode}^{str} = 0.06(0.99)^{episode}$ ¹⁵ $\epsilon_{episode} = \frac{0.5}{\ln(episode + 1)}$
Experimentation Condition	500 episodes trail 50 different trails	10000 episodes trial 100 different trials	150 time steps lift-up $\Delta T = 0.005$ 50 episodes trial 50 different trials ¹⁵ $z^{1,2}(1) \in U(1,2)$ ¹⁶ $z^3(1) \in U(2,3)$	300 time steps lift-up $\Delta T = 0.005$ 50 episodes learning trial (exploration is on) 50 episodes testing trial (exploration is off) 100 different trials $z^{1,2,3}(1) \in U(2,3)$ Structure is updated every three (one) episodes for concurrent beh./str. (str.) learning
Performance Measures	Received reinforcement signal (MA ¹⁷ filtered with window size 10)	1-Received reinforcement signal (MA filtered with window size 100) 2-Probability of Error (calculated every 1000 episodes)	1-Received reinforcement signal 2-Behavioral performance	1-Received reinforcement signal 2-Behavioral performance

جدول ۱-۶. شرایط آزمایش برای آزمون‌های یادگیری

¹⁵ Note that $\alpha_{episode}^{str}$ is not used in the only behavior learning case, and $\alpha_{episode}^{beh}$ and $\epsilon_{episode}$ are not used in the comparison of structure learning alone case. $\epsilon_{episode}$ is set zero in episodes 51-100 (testing episodes).

¹⁶ $U(a, b)$ is a uniform random distribution between a and b.

¹⁷ Moving Average

	Abstract Problem	Object Lifting
Problem Space	5x5 problem space and 7 actions	Multi-robot object lifting
Solution Space	Seven behaviors each can produce has B_i one type of action (behavior) $A_i = \{a_i, NA\}$	Five naïve behaviors' state and action space without any specific mapping (<i>Push more, Don't go fast, Stop at goal, Hurry up, and Slow down.</i>)
Evolution Parameters	Generations = 200, Population size = 30, No. of individual evaluations in each gen. = 300, Tournament selection (competition between 3 individuals), the best individual of each genetic pool goes directly to the $p_c = 0.5$ next generation, $p_m^{hard} = \begin{cases} \frac{0.01}{\log_2(gen+1)} & \text{gen} \leq 179 \\ 0 & 180 \leq \text{gen} \leq 200 \end{cases}$ $p_m^{soft} = 2p_m^{hard}$	Generations = 50, Population size = 20, No. of individual evaluations in each gene. = 200, Tournament selection (competition between 3 individuals), the best individual of each genetic pool goes $p_c = 0.5$ directly to the next generation, $p_m^{hard} = \begin{cases} \frac{0.01}{\log_2(gen+1)} & \text{gen} \leq 40 \\ 0 & 41 \leq \text{gen} \leq 50 \end{cases}$ $p_m^{soft} = 2p_m^{hard}$
Learning Parameters	ZO Structure Learning $\alpha_{episode} = 1/episodes$	ZO Structure Learning $\alpha_{episode}^{str} = 0.1(0.99)^{episode}$
Memetic Parameters	Culture size = 5 $(5-12)\alpha_{T_i} = 0.3$ Meme selection: scaled roulette wheel (p=0.9)/random (p=0.1)	Culture size = 5 $(5-12)\alpha_{T_i} = 0.3$ Meme selection: scaled roulette wheel (p=0.9)/random (p=0.1)
Experimentation Condition	100 episodes trial 3 runs Fitness is the mean value of reinforcement signal during the agent's lifetime	$\Delta T = 0.005$ 300 time step lift-up, 25 (5) episodes trial for with structure learning (fixed structure) case 3 runs for learning cases – 2 runs for fix structure case - 1 run for pre-defined behaviors (no evolution) $z^{1,2,3}(1) \in U(2,3)$ Fitness (Lift time fitness) is calculated by averaging the last 5 episodes (all episodes) - Fitness is used for evolutionary mechanism (and not Life time fitness)
Performance Measures	Received reinforcement signal (MA filtered with window size 10)	1-Received reinforcement signal 2-Behavioral performance

جدول ۲-۶. شرایط آزمایش برای آزمون‌های تکامل/الگوریتم ممیتیک

۷) تحلیل احتمالاتی معماری PPSSA

۷-۱) مقدمه

در فصل‌های پیش به چگونگی یادگیری و تکامل در معماری رفتارگرایی سلسله مراتبی پرداختیم. چگونگی تخصیص امتیاز بین رفتارهای مختلف و روش به روز رسانی دانش عامل نسبت به محیط اطرافش توضیح داده شد. این که عامل تا چه میزان یاد بگیرد و به چه بازدهی برسد البته موضوعی است مجزا که به مساله‌ی مورد نظر، چگونگی انتخاب زیرفضای حالت هر رفتار، سیگنال تقویت طراحی شده و ... بستگی دارد. با این همه می‌توان با فرض گرفتن چند ویژگی، رفتار معماری را تا حدی تحلیل کرد. در این فصل سعی در تحلیل احتمالاتی معماری رفتارگرایی PPSSA خواهیم کرد. در این راستا ابتدا احتمال کنترل‌کننده بودن هر کدام از رفتارهای معماری محاسبه می‌شود و سپس تخمین‌ای از زمان لازم برای یادگیری رفتارها به دست می‌آید. نتایج جالب توجه‌ای در این میان حاصل می‌شود که می‌تواند طراح را در انتخاب به‌تر رفتارها و معماری یاری کند.

مانند بخش‌های پیش، معماری PPSSA را برمی‌گزینیم که هر رفتار B_i در زیرمجموعه‌ی $S_i \subset S$ تحریک می‌شود و می‌آواند عمل‌ای از مجموعه‌ی A_i انتخاب کند. تعاریف تحریک شدن، فعال شدن و کنترل‌کننده بودن مانند پیش است.

برای خوانایی فرض می‌کنم که هر B_i در لایه‌ی i ام (L_i) قرار گرفته است - نام‌گذاری رفتارها با نام‌گذاری لایه‌ها یکسان است. مشخص است که این فرض محدود کننده نیست و تنها برای جلوگیری از پیچیده شدن روابط در نظر گرفته شده است.

۷-۲) احتمال کنترل‌کننده بودن رفتارها

در این بخش به محاسبه‌ی احتمال کنترل‌کننده بودن رفتارها می‌پردازیم. ابتدا از حالت‌های ساده شروع می‌کنم و به تدریج بر پیچیدگی و دقت تحلیل می‌افزایم.

محاسبه‌ی احتمال این که عمل a_i رفتار B_i - که در لایه‌ی L_i قرار دارد- کنترل عامل را در اختیار بگیرد مورد نظر است. با انجام این کار می‌توان فهمید که هر کدام از رفتارها با چه احتمالی در تصمیم‌گیری‌ی سیستم عامل نقش دارند. با استفاده از این دانش می‌توان زمان لازم برای یادگیری‌ی هر کدام از رفتارها را تخمین زد. به این شکل که هر چه یک رفتار بیش‌تر کنترل عامل را در اختیار بگیرد و تجربه‌ی بیش‌تری از تعامل با محیط کسب کند، شانس بیش‌تری برای یادگیری پیدا می‌کند و احتمالاً سرعت هم‌گرایی به پاسخ مطلوب‌اش بیش‌تر می‌شود.

ایده‌ی محاسبه‌ی احتمال در اختیار گرفتن کنترل عامل توسط یک رفتار در لایه‌ای خاص (مثلاً لایه‌ی i ام) که توسط $P\{a = B_i(s) \text{ and } B_i \text{ is controlling}\}$ نمایش‌اش می‌دهم بدین صورت است: این احتمال برابر است با احتمال این که هیچ‌کدام از لایه‌های بالایی لایه‌ی کنترلی نباشد و همچنین این لایه می‌بایست لایه‌ی کنترلی باشد و عمل‌ای که انتخاب می‌کند نیز این عمل مورد نظر باشد. شرط کنترل‌کننده نبودن یک لایه این است که فعال نشود. فعال نشدن معادل یکی از دو روی‌داد زیر است:

۱. رفتار آن لایه تحریک نشود.

۲. رفتار آن لایه تحریک بشود ولی عمل NA را انتخاب کند.

اگر فرض کنیم معماری دارای m لایه است که پایین‌ترین لایه با شماره‌ی ۱ و بالاترین آن با شماره‌ی m مشخص شده باشند، آن‌گاه می‌توان نوشت:

$$P\{a = B_i(s) \text{ and } B_i \text{ is controlling}\} = P\{\text{upper layers are not active in } s\} \times P\{B_i(s) \text{ is excited}\} \times P\{a = B_i(s) | B_i(s) \text{ is excited}\} \quad (7-1)$$

که با بسط جمله‌ی اول داریم:

$$P\{a = B_i(s) \text{ and } B_i \text{ is controlling}\} = \left(\prod_{j=i+1}^m P\{B_j(s) \text{ is not active}\} \right) \times P_i^e(s) \times P\{a = B_i(s) \mid B_i(s) \text{ is excited}\} \quad (7-2)$$

که در آن $P_i^e(s)$ احتمال تحریک شدن رفتار B_i است.

برای محاسبه‌ی این عبارات می‌بایست $P_i^e(s)$ و احتمال فعال شدن هر لایه را بدانیم. دانستن اولی بستگی به مساله دارد. محاسبه‌ی دومی نیز نه تنها به احتمال تحریک شدن بستگی دارد، بلکه به نوع انتخاب عمل و تجربه‌ی عامل در محیط نیز بستگی دارد. با این همه با فرض‌هایی سعی در تحلیل احتمال‌های بالا می‌کنیم.

چگونگی‌ی ارتباط احتمال فعال شدن (یا نشدن) هر رفتار را با توجه به نوع روش انتخاب عمل و احتمال تحریک شدن‌اش بدین‌گونه می‌توان نوشت:

- Random Action Selection (or Uniform Q-Table)

$$P\{B_i(s) \text{ is not active}\} = P\{B_i(s) \text{ is not excited}\} + P\{B_i \text{ is excited and selects } NA\} =$$

$$\circ \quad (1 - P_i^e(s)) + P_i^e(s) \frac{1}{|A_i| + 1}$$

$$(7-3)$$

که جمله‌ی اول بیان‌گر احتمال تحریک نشدن (و در نتیجه، فعال نشدن) است و جمله‌ی دوم بیان‌گر احتمال تحریک شدن و انتخاب NA . به همین ترتیب خواهیم داشت

- ϵ -greedy Action Selection

$$\circ \quad P\{B_i(s) \text{ is not active}\} = (1 - P_i^e(s)) + P_i^e(s) \left[\frac{P\{Q_i(s, NA) > Q_i(s, a'); \forall a' \neq NA\}(1 - \varepsilon)}{1 + |A_i|} + \varepsilon \right] \quad (7-4)$$

• Boltzman Action Selection

$$\circ \quad P\{B_i(s) \text{ is not active}\} = (1 - P_i^e(s)) + P_i^e(s) \frac{\exp\left(\frac{Q_i(s, NA)}{T}\right)}{\sum_{a' \in A_i'} \exp\left(\frac{Q_i(s, a')}{T}\right)} \quad (7-5)$$

مشاهده می‌شود که برای محاسبه‌ی این احتمال هم نیاز به احتمال تحریک شدن وجود دارد و هم اطلاعاتی راجع به دانش عامل ذخیره شده در Q-Table آن.

ابتدا در نظر بگیرید که $S_i = S$ باشد و در نتیجه همه‌ی رفتارها در همه حالت تحریک شده‌اند. در ضمن فرض کنید که بعد فضای عمل‌های هر کدام نیز A_i باشد. آن‌گاه احتمال فعال نشدن برابر با احتمال انتخاب NA است. اگر نوع انتخاب عمل را تصادفی در نظر بگیریم، می‌توان نوشت:

$$P\{B_i \text{ is not active}\} = P\{B_i \text{ is not excited}\} + P\{B_i \text{ is excited and selects } NA\} = (1 - 1) + 1 \frac{1}{|A_i| + 1} = \frac{1}{|A_i| + 1} \quad (7-6)$$

حال به عنوان نمونه اگر فرض کنیم که فضای عمل‌ها بین رفتارها افزار شده باشد و هر کدام از رفتارها مسوول انتخاب تنها یکی از رفتارها هستند (مثلا یکی رفتار بالا بردن جسم و یکی متوقف شدن و دیگری پایین آوردن) آن‌گاه خواهیم داشت:

$$P\{B_i \text{ is not active}\} = \frac{1}{1 + 1} = \frac{1}{2} \quad (7-7)$$

با استفاده از (7-2) می‌توان نوشت:

$$P\{a_{[\in A_i]} = B_i(s) \text{ and } B_i \text{ is controlling}\} = \left(\prod_{j=i+1}^{m=|A|} \frac{1}{2} \right) \times 1 \times \frac{1}{2} = \left(\frac{1}{2} \right)^{|A|-i+1} \quad (7-8)$$

البته توجه کنید که a مشخص شده در عبارت بالا، هر عمل ای نمی تواند باشد و حتما می بایست متعلق به مجموعه ی تک عضوی A_i باشد. در مقایسه با این عبارت، می توان احتمال انتخاب یک عمل خاص در flat Q-Table با انتخاب تصادفی را آورد

$$P\{a = B(s) | \text{flat}\} = \frac{1}{|A|} \quad (7-9)$$

این گونه تقسیم بندی ی فضا - که همه ی رفتارها با هم تحریک می شوند- عملا آن چیزی نیست که در یک PPSSA رخ می دهد. رفتارها تخصصی هستند و سنسورهای متفاوتی را درک می کنند. به همین دلیل، در هر لحظه تنها بخش کوچکی از کل رفتارها فعال می شوند. اگر بخواهیم این نوع نگاه به مساله را به بیان ای ریاضی بنویسیم، می توان احتمال تحریک هر رفتار را بدین گونه نوشت:

$$P_i^e(s) = \frac{l}{m} \quad (7-10)$$

که بدان معناست که در هر لحظه، به طور میانگین l رفتار از میان m رفتار تحریک شده اند. با چنین فرضی و با در نظر گرفتن انتخاب عمل تصادفی می توان نوشت:

$$P\{a = B_i(s) \text{ and } B_i \text{ is controlling}\} = \left[\prod_{j=i+1}^m \left(1 - \frac{l}{m} + \frac{l}{m} \frac{1}{|A_i|+1} \right) \right] \times \frac{l}{m} \times \frac{1}{|A_i|+1} \quad (7-11)$$

که اگر مانند قبل عمل ها را بین رفتارها افراز تک تک کرده باشیم - یعنی $m = |A|$ - خواهیم داشت

$$\left. \begin{array}{l} |A_i| = 1 \\ \bigcup_{i=1}^{|A|} A_i \cap A_j = \emptyset; \forall i, j \wedge i \neq j \end{array} \right\} \Rightarrow P\{a = B_i(s) \text{ and } B_i \text{ is controlling}\} = \frac{l}{2m} \left(1 - \frac{l}{2m}\right)^{m-i} \quad (7-12)$$

که اگر در آن تعداد کمی از رفتارها به طور هم‌زمان تحریک شوند، آن‌گاه

$$\frac{l}{m} \ll 1 \Rightarrow P\{a = B_i(s) \text{ and } B_i(s) \text{ is controlling}\} \approx \frac{l}{2m} \left(1 - \frac{l(m-i)}{2m}\right) \quad (7-13)$$

حال فرض پیشین در مورد احتمال تحریک لایه‌ها را کمی تعمیم می‌بخشیم: احتمال تحریک لایه‌ها را برابر نمی‌گیریم. می‌توان گفت که در SSA، فرض بر این است که لایه‌های پایین‌تر لایه‌هایی هستند که در اکثر مواقع تحریک شده‌اند و کنترل عامل را بر عهده دارند و در عوض لایه‌های بالایی آن‌هایی‌اند که به ندرت - و تنها در شرایط بحرانی - فعال می‌شوند. این فرض معقول را می‌توان بدین‌گونه بیان کرد:

$$P_i^e(s) > P_{i+1}^e(s) \quad i = 1, \dots, m-1 \quad (7-14)$$

دو مدل ساده برای تغییرات P_i^e در نظر می‌گیریم. یکی مدل‌ایست که احتمال‌های مورد بحث به صورت هندسی از پایین به بالا کاهش یابند و دیگر مدل‌ای است که آن احتمال‌ها به طور حسابی از پایین به بالا کاهش می‌یابند. به عبارت دقیق‌تر برای مدل کاهش حسابی داریم:

$$P_{i+1}^e = \eta P_i^e \quad (7-15)$$

و برای مدل حسابی داریم:

$$P_{i+1}^e = \max(P_i^e - \delta, 0) \quad (7-16)$$

محاسبات احتمال کنترل کننده بودن را برای مدل کاهش هندسی انجام می دهیم: به مانند پیش انتخاب عمل مان را تصادفی فرض می کنیم. هم چنین $A_i = A$ در نظر می گیریم. خواهیم داشت:

$$P_{i+1}^e = \eta P_i^e \Rightarrow P_i^e = \eta^{i-1} P_1^e \quad (7-17)$$

$$P\{a = B_i(s) \mid B_i(s) \text{ is exited}\} = \frac{1}{|A|+1} \quad (7-18)$$

$$P\{B_i(s) \text{ is not active}\} = (1 - P_i^e(s)) + P_i^e(s) \frac{1}{|A|+1} = (1 - \eta^{i-1} P_1^e) + \frac{\eta^{i-1} P_1^e}{|A|+1} = \frac{1}{|A|+1} [1 + |A|(1 - \eta^{i-1} P_1^e)] \quad (7-19)$$

پس در نتیجه داریم:

$$P\{a = B_i(s) \text{ and } B_i(s) \text{ is controlling}\} = \frac{1}{(|A|+1)^{m-i}} \prod_{j=i+1}^m (1 + |A|(1 - \eta^{j-1} P_1^e)) \cdot \frac{\eta^{i-1} P_1^e}{|A|+1} \quad (7-20)$$

از تحلیل تحلیل ریاضی فرم نزول حسابی احتمال تحریک صرف نظر می شود و تنها نتایج شبیه سازی شده ی آن آورده می شود.

در شبیه سازی ها دیده می شود که پارامترهای مختلف مانند تعداد لایه ها (m)، بعد فضای خروجی (A) و مخصوصا نرخ کاهش احتمال تحریک (η و ε) می توانند تاثیری قابل توجه بر احتمال کنترل کننده بودن لایه ها بگذارند. نتایج این شبیه سازی ها در نمودارهای بعدی می آید.

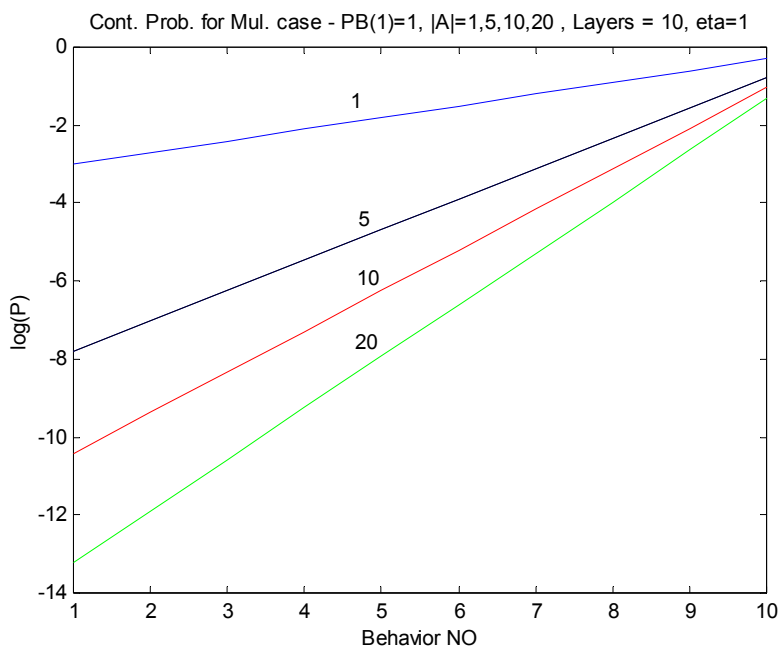
ابتدا فرض می‌کنیم هیچ کاهش احتمال تحریک‌ای وجود نداشته باشد ($\eta = 1$) و تنها بعد فضای خروجی تغییر کند. از رابطه‌ی به دست آمده مشاهده می‌شود که هر چه این بعد بزرگ‌تر باشد، احتمال کنترل‌کننده بودن برای لایه‌های پایین‌تر کمتر می‌شود. این اثر را در شکل ۷-۱ می‌بینید. همچنین تاثیر تعداد لایه‌ها نیز در شکل ۷-۲ بررسی شده است. مشاهده می‌شود که هر چه تعداد لایه‌های بالای یک رفتار بیش‌تر باشد، احتمال کنترل‌کننده بودن آن کمتر می‌شود و نرخ کاهش آن نیز نمایی است. این دو نتیجه علاوه بر آن که به راحتی از روی رابطه‌ی قبل قابل حصول‌اند، از نظر شهودی نیز پذیرفتی‌اند.

اینک تاثیر η را مطالعه می‌کنم. η میزان کاهش احتمال تحریک‌شدن در مدل هندسی را مشخص می‌کند و هر چقدر اندازه‌اش کم‌تر باشد، شانس تحریک لایه‌های بالایی کمتر می‌شود. تاثیر این پارامتر را در شکل‌های ۷-۳ و ۷-۴ می‌بینید. مشاهده می‌شود که هنگامی که η از مقدار ۱ کم‌تر می‌شود، شیب کاهش نمودار احتمال کنترل‌کننده بودن کم می‌شود و به حالت تخت نزدیک می‌شود. این به آن معنا است که احتمال کنترل‌کننده بودن همه‌ی لایه‌ها کم و بیش یک‌سان می‌شود. با کاهش بیش‌تر η ، حتی جهت این شیب نیز تغییر می‌کند و احتمال کنترل‌کننده بودن لایه‌های بالا کم‌تر از لایه‌های پایین می‌شود.

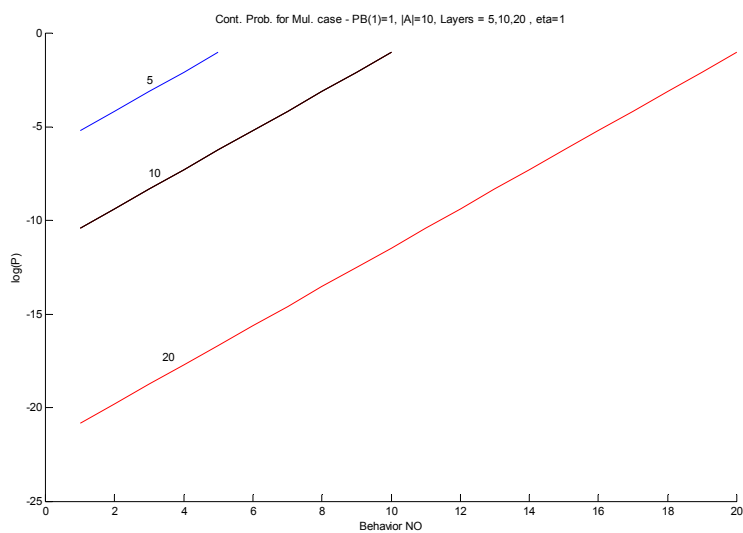
اینک مدل کاهش احتمال تحریک حسابی را مورد بررسی قرار می‌دهیم. در این مدل - همان‌طور که پیش‌تر گفته شد- احتمال تحریک پایین‌ترین لایه بیش‌تر از بقیه است و هر چه به لایه‌های بالا می‌رویم، این احتمال به اندازه‌ی ثابتی کاهش می‌یابد. مانند قبل حالت‌های مختلفی را بررسی می‌کنم. اگر $\varepsilon = 0.1$ و تعداد لایه‌ها ۱۰ باشد، آن‌گاه تاثیر تغییر بعد فضای عمل‌ها (A) در شکل ۷-۵ نمایش داده شده است. همان‌طور که انتظار می‌رفت، با افزایش بعد، احتمال کنترل‌کننده بودن لایه‌ها کاهش می‌یابد. لازم است توجه کنید که اگر $\varepsilon = 0$ می‌بود، آن‌گاه نمودارهایی دقیقاً معادل شکل ۷-۱ به دست می‌آمد. در ضمن می‌بایست توجه شود که به علت غیر صفر بودن ε ، منحنی‌ها در مقیاس لگاریتمی دارای انحنا هستند و کاملاً با خط راست تطابق ندارند (خط راست معادل کاهش نمایی احتمال است). با توجه به شکل منحنی، وجود انحنا معادل نرخ کاهش کم‌تر یا بیش‌تر از نمایی ثابت است). با بررسی تاثیر تعداد لایه‌ها بر احتمال کنترل‌کننده بودن مشخص می‌شود که به مانند

قبل هر چه تعداد لایه‌ها بیش‌تر باشد، در مقادیر کوچک ε (مثلاً $\varepsilon = 0.05$ در این مورد)، آن احتمال برای لایه‌های پایین‌تر کم‌تر می‌شود (شکل ۶-۷). البته این به شرطی است که مقدار ε به گونه‌ای تنظیم شده باشد که با پایین آمدن در لایه‌ها، احتمال نیز کم شود. دیده می‌شود که به ازای مقادیر بزرگ ε ، این جهت معکوس می‌تواند باشد.

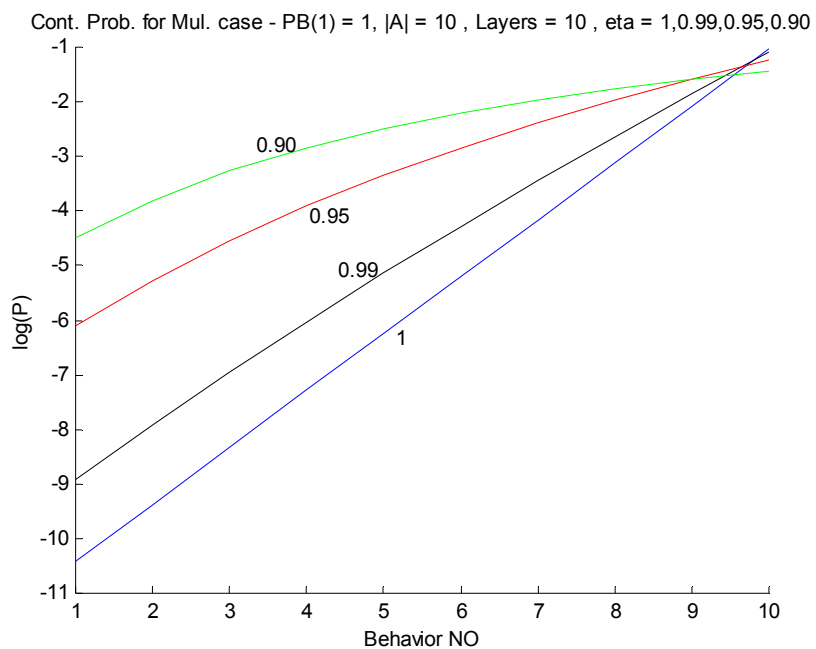
در نهایت تاثیر ε را در شکل ۷-۷ مشاهده می‌کنید. مطابق قبل می‌بینیم که با افزایش میزان کاهش، منحنی‌ها تخت‌تر می‌شوند و در نتیجه احتمال کنترل‌کننده بودن رفتار بیش‌تر می‌شود. توجه کنید که مقدار ε بیش از حدی نمی‌تواند باشد چون معادل احتمال تحریک شدن صفر برای لایه‌های بالا خواهد بود.



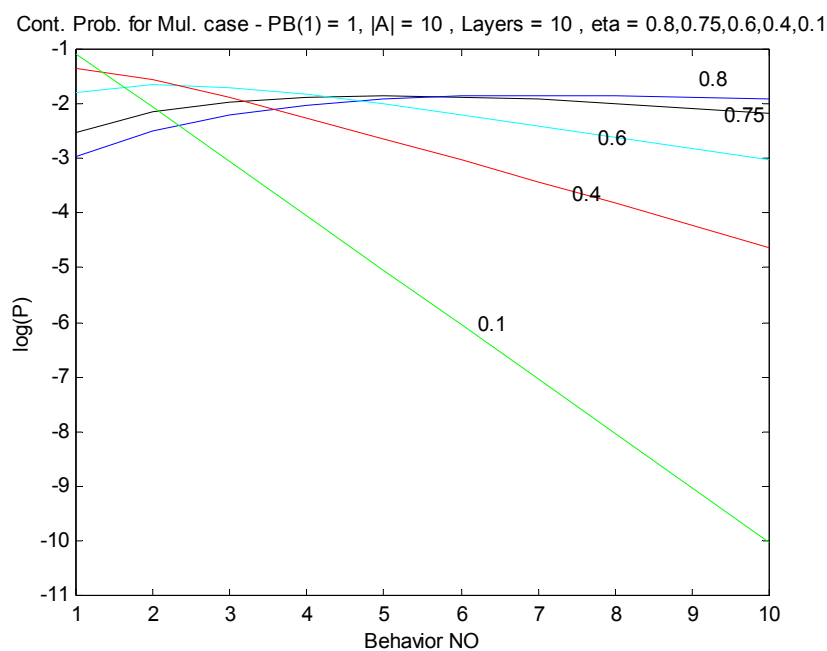
شکل ۷-۱. تاثیر بعد خروجی بر احتمال کنترل کننده بودن. هر چه بعد خروجی بزرگتر باشد، این احتمال کم‌تر می‌شود.



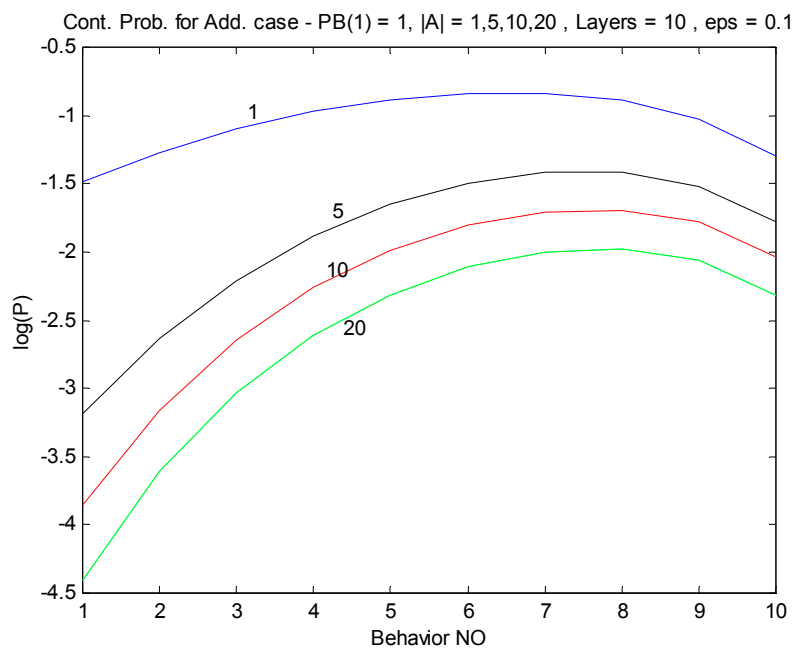
شکل ۷-۲. تاثیر تعداد رفتارها بر احتمال کنترل کننده بودن. هر چه تعداد رفتارها بیش‌تر باشد، رفتارهای پایین دارای احتمال کنترل کننده بودن کم‌تری هستند.



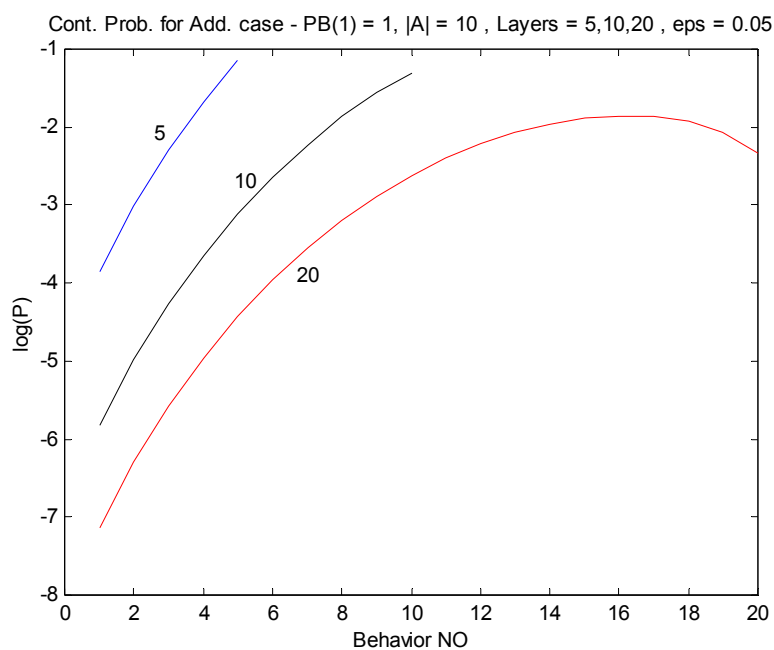
شکل ۷-۳. تاثیر η بر احتمال کنترل کننده بودن. در مقادیر نزدیک به ۱ η ، با کاهش مقدار نمودارها تخت تر می شوند.



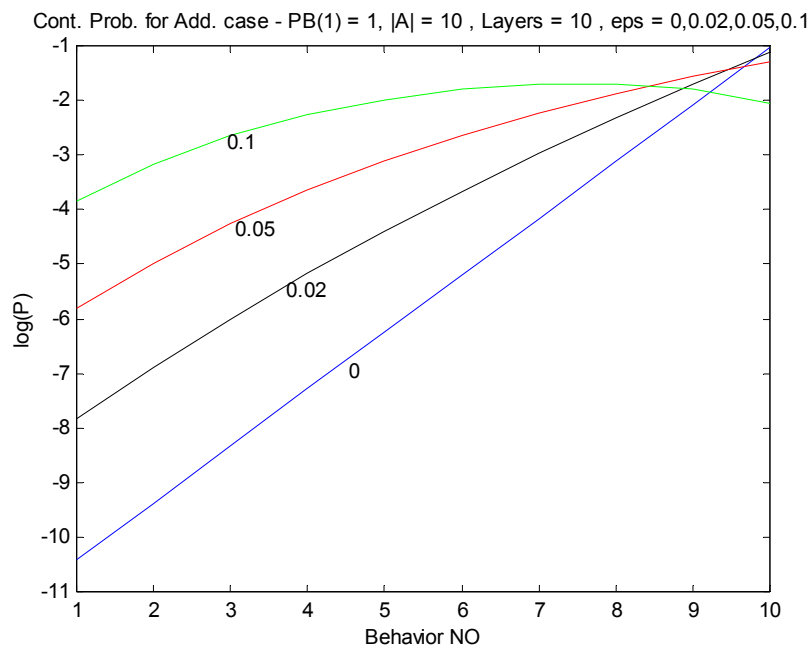
شکل ۷-۴. تاثیر η بر احتمال کنترل کننده بودن. مقادیر کوچک η باعث افزایش شدید احتمال کنترل کننده بودن لایه های پایین و کم شدن آن احتمال برای لایه های بالا می شود. با کاهش مقدار نمودارها تخت تر می شوند.



شکل ۵-۷. تأثیر بعد خروجی بر احتمال کنترل کننده بودن در مدل کاهش حسابی. هر چقدر بعد خروجی بزرگتر باشد، این احتمال کمتر می شود.



شکل ۶-۷. تأثیر تعداد رفتارها بر احتمال کنترل کننده بودن. هر چه تعداد رفتارها بیش تر باشد، رفتارهای پایین دارای احتمال کنترل کننده بودن کمتری هستند.



شکل ۷-۷. تاثیر ϵ بر احتمال کنترل کننده بودن. با افزایش ϵ ، منحنی‌ها تخت‌تر می‌شوند.

۷-۳) تحلیل‌ای بر سرعت یادگیری در معماری PPSSA

هدف از تحلیل‌های احتمالاتی PPSSA بخش پیش، درک بیش‌تر و عمیق‌تر از نحوه‌ی کارکرد عامل رفتارگرای دارای معماری PPSSA در محیط است. در قسمت قبل احتمال کنترل‌کننده بودن هر کدام از رفتارها در معماری را برای چند حالت مختلف به دست آوردیم. در این بخش با استفاده از آن اطلاعات، تخمین‌ای از زمان لازم برای یادگیری‌ی هر کدام از آن رفتارها را محاسبه می‌کنیم. برای مشخص شدن نحوه‌ی کار فرض کنید که معماری‌ی با ترتیب رفتارهای ثابت‌ای داریم که قرار است با استفاده از روش یادگیری تقویتی –همان‌طور که در فصل‌های پیش توضیح دادیم– نحوه‌ی عمل‌کرد بهینه‌ی هر کدام از رفتارها را بیابد (مساله‌ی یادگیری رفتار). به طور تقریبی می‌توان گفت که هر چه یک رفتار (دقیق‌تر: حالت-عمل یک رفتار) بیش‌تر کنترل‌کننده شود و سیگنال تقویت دریافت کند آن‌گاه به احتمال زیاد تخمین دقیق‌تری نسبت به ارزش خود خواهد داشت و می‌توان ادعا کرد که آن رفتار بیش‌تر از رفتارهایی با تعداد تجربه‌ی کم‌تر یاد گرفته است. در این بخش به عنوان معیاری برای یادگیری تعداد تجربه‌های کنترل‌کنندگی (r) مورد نظر است. در زیربخش‌های بعدی توزیع احتمال تعداد تجربه‌های کنترل‌کنندگی را با تعمیم توزیع احتمال دو جمله‌ای منفی به دست می‌آوریم.

۷-۳-۱) توزیع دو جمله‌ای منفی

توزیع دو جمله‌ای منفی، بیان‌گر تعداد آزمایش لازم برای r بار رخ دادن پیش‌آمدی X با احتمال p است. چگالی احتمال، متوسط و واریانس آن بدین‌گونه است:

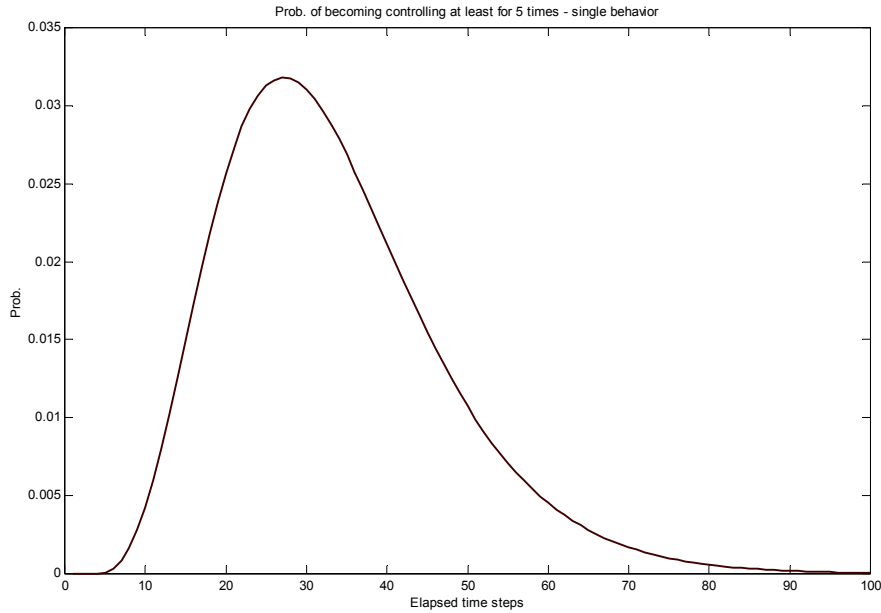
$$P_X(i) = \begin{cases} \binom{i-1}{r-1} p^r (1-p)^{i-r} & i = r, r+1, \dots \\ 0 & \text{otherwise} \end{cases} \quad (7-21)$$

$$m_X = \frac{r}{p} \quad (7-22)$$

$$\sigma_X^2 = \frac{r(1-p)}{p^2} \quad (7-23)$$

فرض کنید معماری ای تک-رفتاری داریم و احتمال کنترل کننده بودن آن تک-رفتار، P_{B_i} باشد. حال برای محاسبه ی چگالی ی احتمال تعداد آزمایش لازم برای r بار کنترل کننده بودن آن، می توان از رابطه ی چگالی احتمال توزیع دو جمله ای منفی استفاده کرد. در نتیجه می توان دانست به طور احتمالی چه میزان تجربه برای یادگیری ی این رفتار لازم است. به عنوان نمونه در شکل ۷-۸، این چگالی ی احتمال برای ۵ بار کنترل کننده بودن تک-رفتاری با احتمال کنترل کننده بودن ۰,۱۵، نمایش داده شده است. بدیهی ست که برای مقادیر کمتر از ۵، چگالی صفر باشد. به راحتی به دست می آید که متوسط زمان لازم، ۳۳,۳۳ واحد زمانی است.

حال این ایده را برای یک معماری ی کامل گسترش می دهیم. در این جا ما به جای یک احتمال، مجموعه ای از احتمال های کنترل کننده بودن (P_{B_i}) را داریم. فرآیندهای این احتمال ها - احتمال کنترل کننده بودن یک رفتار- مستقل نیستند و رخ داد یکی شان به دیگری وابسته است. البته مشخص است که در هر لحظه فقط یکی از این فرآیندها (کنترل کنندگی ی یک رفتار) رخ می دهد. برای به دست آوردن توزیع احتمال این که همه ی آن رفتارها به میزان r بار کنترل کننده باشند می توان ایده ی توزیع دو جمله ای منفی را گسترش داد. اما پیش از انجام آن کار، تقریبی حدودی از این زمان به دست می آوریم.



شکل ۸-۷. چگالی احتمال تعداد تلاش لازم برای ۵ بار کنترل کننده بودن رفتاری با احتمال کنترل کننده بودن ۰,۱۵ در هر واحد زمانی.

فرض کنید که $X_1, X_2, \dots, X_m$ فرآیند r بار کنترل کننده بودن هر کدام از رفتارهای $B_1, B_2, \dots, B_m$ معماری اند در صورتی که فرض کنیم آن‌ها درون معماری نیستند و فقط احتمال کنترل کننده بودن شان جوری تنظیم شده است که P_{B_i} باشد^۱. در نتیجه رخداد هر کدام شان مستقل از بقیه اند - زیرا فرض شده بقیه ای وجود ندارد. پس می توان گفت هر کدام از این ها دارای توزیع دو جمله ای منفی با پارامترهای r و P_{B_i} هستند. توجه کنید که در این آزمایش شرط مستقل بودن به صورت اجباری لحاظ شده است و در نتیجه چنین فرآیندی معادل فرآیند اصلی - که رفتارها در معماری قرار دارند - نیست.

اگر توزیع زمان لازم برای r بار کنترل کننده بودن در معماری ی واقعی را با $Y_1, Y_2, \dots, Y_m$ نشان دهیم، خواهیم داشت:

$$m_{Y_i} \geq m_{X_i} \quad i = 1, \dots, m \quad (7-24)$$

^۱ تعبیر دیگری از این تعریف فرآیند به زودی خواهد آمد.

دلیل این موضوع این است: هنگامی که رفتار B_i در معماری قرار ندارد، به احتمال P_{B_i} فعال می‌شود و به شمارنده‌ی تعداد فعالیت‌های اش یکی اضافه شود. این احتمال در هر کدام از واحدهای زمانی قابل اعمال است و انتخاب آزادانه است. اما وقتی این رفتار در معماری قرار گرفت، مسلم است که تعداد واحدهای زمانی‌ای که در اختیار دارد بیش‌تر نیست – چون بعضی از آن واحدهای زمانی به دیگر رفتارها اختصاص داده شده است. پس امید ریاضی زمان لازم برای r بار تکرار حداقل به اندازه‌ی امید ریاضی حالت اول است. به عبارت دیگر این قید برای رفتار درون معماری وجود دارد که در صورتی که یکی از رفتارهای دیگر کنترل‌کننده باشد، آن رفتار دیگر کنترل‌کننده نیست.

اگر متوسط زمان لازم را با $\bar{T}$ نشان دهیم، رابطه‌های زیر صادق‌اند:

$$\bar{T} \geq \max_i m_{X_i} \geq \min_i m_{X_i} \quad (7-25)$$

فرآیند دیگری را در نظر می‌گیریم: فرض کنید که اجازه‌ی کنترل‌کننده بودن را به این نحو تقسیم کرده‌ایم که ابتدا می‌گذاریم B_1 ، r بار کنترل‌کننده باشد. یعنی در این حالت اگر رفتار دیگری کنترل‌کننده بود، آن را به عنوان تجربه‌ای برای اش در نظر نمی‌گیریم و تنها تجربه‌های B_1 را می‌شماریم. این نوع اجازه‌ی فعالیت دادن، توزیع احتمال کنترل‌کننده بودن B_1 را عوض نمی‌کند (B_1 هم‌چنان دارای احتمال کنترل‌کننده بودن P_{B_1} خواهد بود) ولی تجربه یادگیری آن را مستقل از بقیه می‌کند. این نوع تقسیم‌بندی، معادل فرآیند X_1 است.² چنین فرآیندی به طور متوسط به m_{X_1} واحد زمانی نیاز دارد. پس از انجام چنین چیزی، همین عمل را برای B_2 ، B_3 و ... تا B_m انجام می‌دهیم. با این کار می‌توان گفت زمان به بخش‌هایی تقسیم شده است که هر بخش تنها به یکی از رفتارها تخصیص داده شده است. طبیعی است که چون تجربه‌های رفتارهای دیگر را هنگامی که در بخش

² تعبیر دیگر اشاره شده در قسمت پیش.

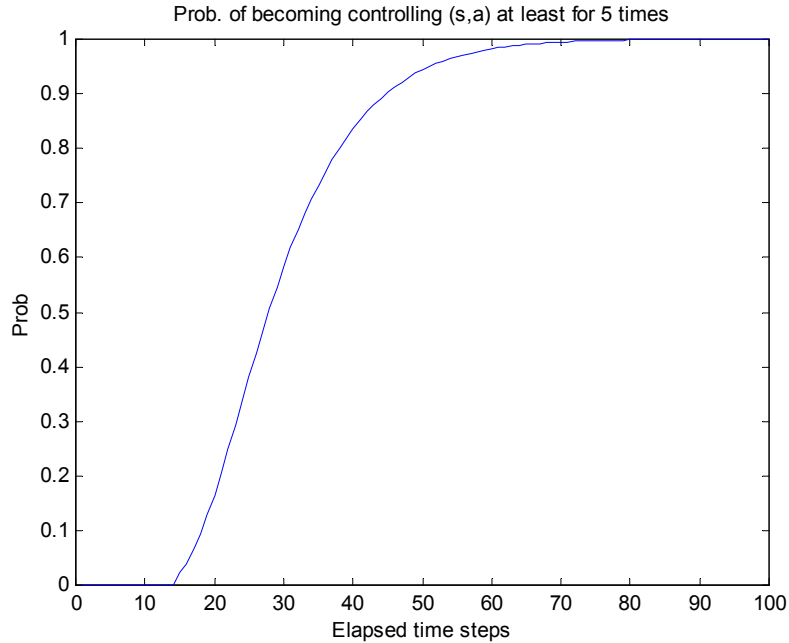
زمانی رفتار خاصی بوده‌ایم نشمارده‌ایم، پس زمان واقعی لازم (در صورتی که معماری مثل شرایط معمول اجرا می‌شد)، کم‌تر خواهد بود. پس خواهیم داشت:

$$\bar{T} \leq m_{X_1} + m_{X_2} + \dots + m_{X_m} \quad (7-26)$$

بدین نحو می‌توان تخمین‌ای از میزان زمان لازم برای یادگیری به دست آورد. مثلاً با توجه به روابط به دست آمده، می‌دانیم که کندترین رفتار (آن‌ای که دارای کوچک‌ترین P_{B_i} است) بخش عمده‌ی زمان لازم برای یادگیری را مشخص می‌کند و یا یادگیری کل معماری از مرتبه‌ی یادگیری هر کدام از رفتارهایی است که با احتمال مشخص شده تجربه می‌کنند. پس می‌توان نتیجه گرفت که اگر لازم است همه‌ی رفتارها یاد گرفته شوند، به‌تر است احتمال کنترل‌کننده بودن‌شان بیش از اندازه کوچک نباشد. البته SSA معمولاً آن‌قدر مقاوم هست که حتی اگر یکی از رفتارها نیز کاملاً یاد گرفته نشده باشد، باز هم رفتار مطلوب از کل معماری حاصل شود. اگرچه باید توجه شود که این گزاره‌ی آخر، الزاماً درست نیست و گزاره‌ای تجربی است تا ریاضی. به طور خلاصه خواهیم داشت:

$$\max_i m_{X_i} \leq \bar{T} \leq m_{X_1} + m_{X_2} + \dots + m_{X_m} \quad (7-27)$$

پس از تحلیل کران بالا و پایین امید ریاضی زمان r تجربه داشتن برای هر رفتار، می‌توان تحلیل دقیق‌تری انجام داد و توزیع احتمال این فرآیند را به دست آورد. برای این کار ایده‌ی توزیع دوجمله‌ای منفی را تعمیم می‌دهم. بدین‌گونه که فرض می‌کنم همه‌ی رفتارها در معماری هستند و هر کدام احتمال کنترل‌کننده بودن مشخصی دارد و حال می‌خواهم احتمال حداقل r بار کنترل‌کننده بودن آن‌ها را محاسبه کنم. توجه کنید که به دلیل نرخ متفاوت کنترل‌کننده بودن، بعضی‌ها از رفتارها بیش‌تر کنترل‌کننده می‌شوند (و در نتیجه بارهای بیش‌تری سیگنال تقویت دریافت می‌کنند) و بعضی کم‌تر. شرط کافی برای یادگیری را این فرض می‌گیریم که همه‌ی آن‌ها حداقل r بار کنترل‌کننده بوده باشند.



شکل ۹-۷. احتمال بیش از ۵ بار کنترل کننده بودن برای معماری دو لایه

ابتدا از تعمیم توزیع دوجمله‌ای به بیش از یک متغیر تصادفی استفاده می‌کنیم. اگر $Y_1, Y_2, \dots$ و Y_m بیان‌گر متغیرهای تصادفی با احتمال $p_1, p_2, \dots, p_m$ و $(p_1 + p_2 + \dots + p_m \leq 1)$ باشند که در هر واحد زمانی تنها یکی از آن‌ها رخ می‌دهد، آن‌گاه احتمال رخداد $r_1, r_2, \dots, r_m$ هر کدامشان در r_0 بار تلاش $r_1 + r_2 + \dots + r_m \leq r_0$ ، بدین گونه خواهد بود:^۳

$$P_{PD}\{Y_1 = r_1, Y_2 = r_2, \dots, Y_m = r_m | r_0\} = \frac{r_0!}{\left(\prod_{i=1}^m r_i!\right) \left(r_0 - \sum_{i=1}^m r_i\right)!} \prod_{i=1}^m p_i^{r_i} \left(1 - \sum_{i=1}^m p_i\right)^{\left(r_0 - \sum_{i=1}^m r_i\right)} \quad (7-28)$$

که در صورتی که شرط‌های مساله ارضا نشود (یعنی $r_1 + r_2 + \dots + r_m \geq r_0$)، این احتمال صفر خواهد بود.

^۳ چنین توزیعی را توزیع چندجمله‌ای می‌نامیم.

خواست مورد نظر - یعنی محاسبه بیش از r بار کنترل کننده بودن در r_0 تلاش - بدین گونه خواهد بود:

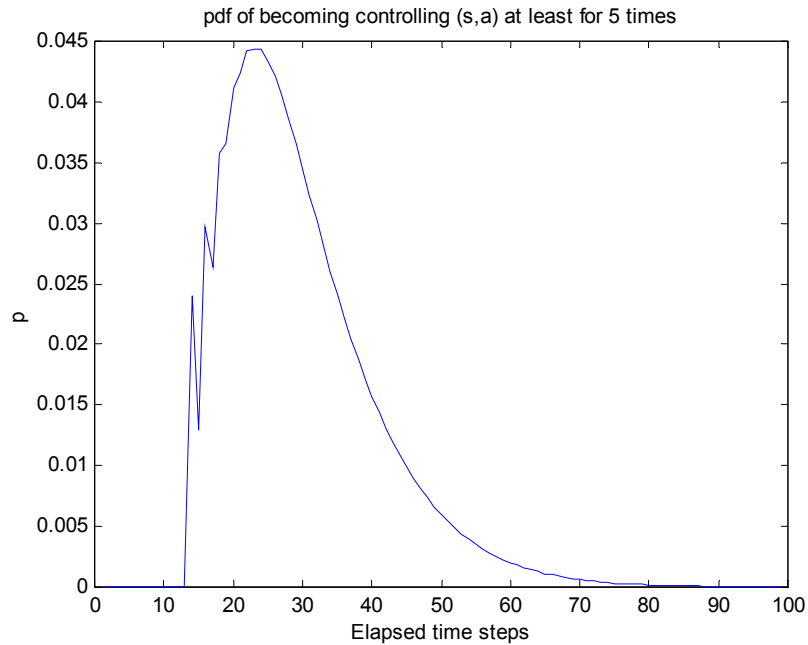
$$P_{GNB}\{(Y_1, Y_2, \dots, Y_m) \geq r | r_0\} = \sum_{r_1=r}^{r_0-(m-1)} \sum_{r_2=r}^{r_0-(m-1)} \dots \sum_{r_m=r}^{r_0-(m-1)} P_{NB}(r_1, r_2, \dots, r_m | r_0) \quad (7-29)$$

به نظر نمی آید فرم بسته ای برای محاسبه ای این احتمال وجود داشته باشد. برای چنین محاسبه ای می بایست از روش های عددی استفاده کرد. با این وجود، استفاده از روش های عددی نیز بسیار ساده نیست چون با افزایش تعداد متغیرهای تصادفی، میزان محاسبات لازم (و در نتیجه زمان آن) به طور نمایی افزایش می یابد.

برای چنین محاسبه ای ابتدا از روش های تخمینی مانند تخمین مونت کارلو استفاده شد، اما پاسخ مناسب و دقیقی به دست نیامد. برای همین همان عبارت بالا به طور مستقیم محاسبه شد با این تفاوت که به جای یک پله بالا رفتن هر کدام از حلقه های جمع ($\sum \dots$)، از پله های چندتایی استفاده شده است. در نتیجه دقت محاسبات اندکی پایین آمده است ولی در عوض سرعت به مقدار قابل توجهی بالا رفته است.

چگونگی استفاده از رابطه ای بالا برای حل مساله ای زمان لازم برای یادگیری رفتارها واضح است. فرض کنیم که می خواهیم احتمال این را که هر کدام از رفتارها r بار یا بیش تر در r_0 تلاش کنترل کننده شوند به دست آوریم. پس در رابطه های پیشین $P_i = P_{B_i} = P\{B_i \text{ is controlling}\}$ را - که در بخش های پیشین برای چند حالت مختلف به دست آمد - قرار می دهیم.

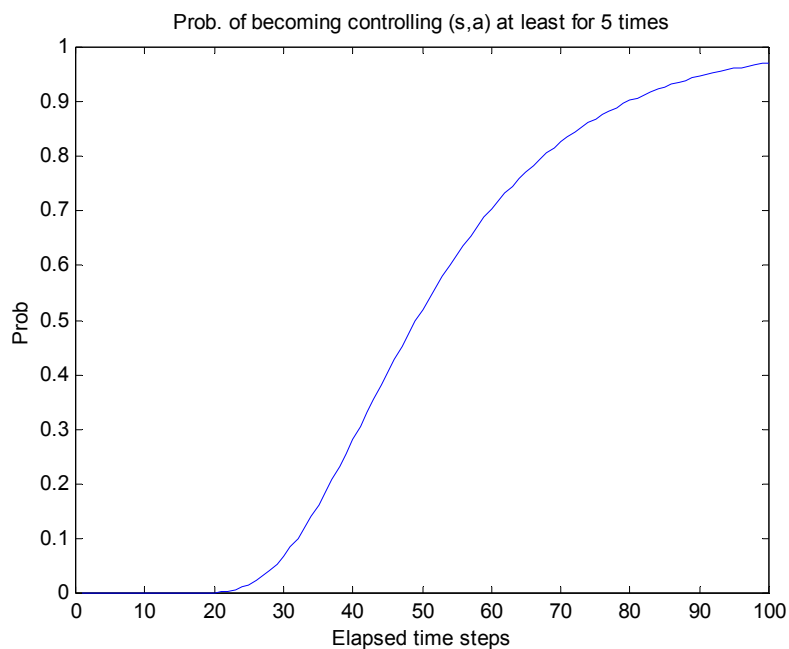
به عنوان مثال، فرض کنید که معماری مان دو لایه است و هر کدام از لایه ها دارای فضای عمل با بعد دو می باشد ($|A|=2$). احتمال تحریک شدن پایین ترین لایه را یک می گیریم ($P_1^e=1$) و ضریب کاهش احتمال را نیز $\eta=0.8$ (به عبارت دیگر $P_2^e=0.8$). در آن صورت، احتمال کنترل کننده بودن این گونه خواهد شد:



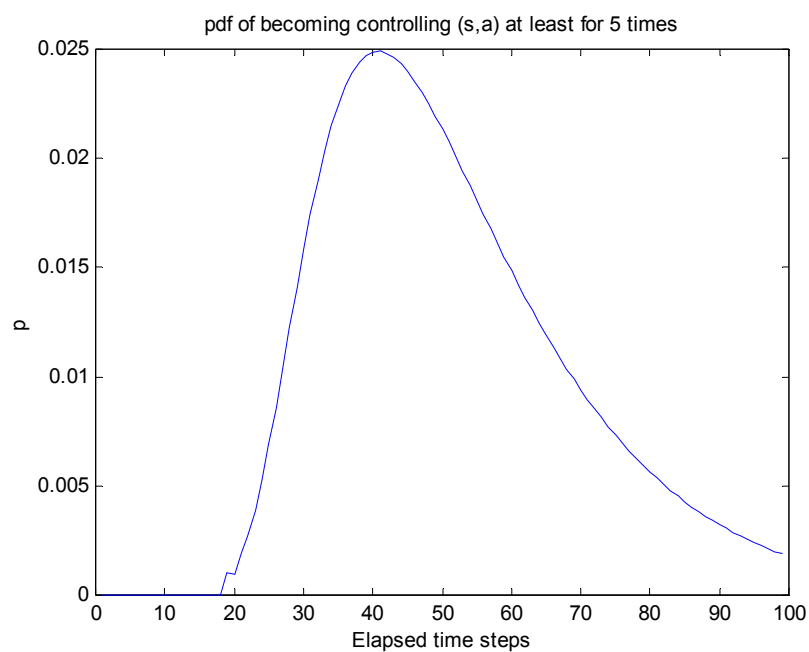
شکل ۷-۱۰. چگالی احتمال بیش از ۵ بار کنترل کننده بودن برای معماری دو لایه

$$\begin{aligned} P_{B_1} &= 0.1556 \\ P_{B_2} &= 0.2667 \end{aligned} \quad (۷-۳۰)$$

$r = 5$ می‌گیریم و r_0 را تغییر می‌دهیم. نتیجه را در شکل ۷-۹ می‌بینید. اگر چگالی توزیع (بر حسب r_0) را رسم کنیم، شکل ۷-۱۰ را خواهیم داشت. قله نمودار در حوالی ۲۵-۳۰ است. با محاسبه‌ی متوسط زمان لازم برای حداقل ۵ بار کنترل کننده بودن، عدد 29.5 به دست می‌آید. این عدد با آنچه از توزیع دوجمله‌ای منفی هر کدام از رفتارها به تنهایی حاصل می‌شود (که در آن‌ها دیگر در معماری نخواهند بود)، نزدیک است، چون:



شکل ۷-۱۱. احتمال بیش از ۵ بار کنترل کننده بودن برای معماری سه لایه



شکل ۷-۱۲. چگالی احتمال بیش از ۵ بار کنترل کننده بودن برای معماری دو لایه

$$P_{B_1} = 0.1556 \Rightarrow m_{B_1} = \frac{5}{0.1556} \approx 32 \quad (7-31)$$

$$P_{B_2} = 0.2667 \Rightarrow m_{B_2} = \frac{5}{0.2667} \approx 29 \quad (7-32)$$

$$m_{T=[B_1 B_2]} = 29.5 \quad (7-33)$$

می‌بینیم که مقدار این دو سری عدد تقریباً نزدیک به هم است.

به عنوان نمونه‌ای دیگر، معماری‌ای سه لایه با همان بعد فضای عمل و $\eta = 0.8$ در نظر بگیرید. در این شرایط، احتمال کنترل‌کننده بودن بدین‌گونه است^۴:

$$\begin{aligned} P_{B_1} &= 0.0892 \\ P_{B_2} &= 0.1529 \\ P_{B_3} &= 0.2133 \end{aligned} \quad (7-34)$$

در این حالت نیز توزیع احتمال و چگالی توزیع احتمال را در شکل‌های ۷-۱۱ و ۷-۱۲ می‌بینید. اگر متوسط زمان لازم برای حداقل ۵ بار کنترل‌کننده بودن را برای کل معماری و همچنین هر کدام از رفتارها حساب کنیم، خواهیم داشت:

$$P_{B_1} = 0.0892 \Rightarrow m_{B_1} = \frac{5}{0.0892} \approx 56 \quad (7-35)$$

$$P_{B_2} = 0.1529 \Rightarrow m_{B_2} = \frac{5}{0.1529} \approx 32 \quad (7-36)$$

$$P_{B_3} = 0.2133 \Rightarrow m_{B_3} = \frac{5}{0.2133} \approx 23 \quad (7-37)$$

$$m_{T=[B_1 B_2 B_3]} \approx 49.3 \quad (7-38)$$

^۴ دلیل تفاوت احتمال کنترل‌کننده بودن دو لایه‌ی پایین با حالت پیش‌با وجود یک‌سان بودن پارامترها- عمل سرکوب لایه‌های بالایی‌ست.

مشاهده می‌شود که کندترین رفتار، زمان کلی را به سمت خود می‌کشاند و زمان کلی لازم - یعنی ۴۹,۳- به ۵۶ پایین‌ترین رفتار بسیار نزدیک است.

۸) تاثیر عدم قطعیت سیگنال تقویت بر تابع ارزش و سیاست عامل

۸-۱) مقدمه

بخش بزرگی از این پایان نامه درباره‌ی استفاده از روش‌های یادگیری تقویتی برای طراحی‌ی خودکار معماری‌ی سلسله مراتبی‌ی رفتارگرا بود. برای طراحی‌ی یک سیستم مبتنی بر یادگیری تقویتی نکاتی چند از جمله چگونگی‌ی انتخاب متغیرهای حالت، چگونگی‌ی بازنمایی دانش، ساختار سلسله مراتب و طراحی‌ی سیگنال تقویت مناسب می‌بایست مورد توجه قرار گیرد. حل هر کدام از این مسایل ساده نیست و پژوهش‌های بسیاری برای پاسخ دادن به آن‌ها انجام شده است و همچنان ادامه دارد. در این فصل به بررسی‌ی بخش کوچک‌ای از یکی از این مسایل -یعنی طراحی سیگنال تقویت- می‌پردازم و تاثیر عدم قطعیت در سیگنال تقویت را بر تابع ارزش و همچنین سیاست عامل به صورت کران‌هایی به دست خواهیم آورد.^۱

یکی از اصول فلسفی‌ی پس طراحی رفتار با استفاده از یادگیری تقویتی، مقاوم بودن پاداش و جزا نسبت به تغییرات محیطی است و در نتیجه توصیف سیگنال تقویت‌ای یک وظیفه مقاوم‌تر و قابل انتقال‌تر از توصیف سیاست انجام آن است. اگر دینامیک محیط یا عامل تغییر کند، سیاست از پیش مشخص شده دیگر مفید نخواهد بود اما عامل یادگیر در صورت دسترسی به سیگنال تقویت مناسب‌ای که بر اعمالش ارزش‌گذاری کند می‌تواند همچنان کارکرد خود را حفظ کند. با این وجود طراحی‌ی سیگنال تقویت‌ای که هدف طراح را برآورد کند چندان ساده نیست. جدا از مسایل‌ای که هدف بسیار مشخص است (مانند ماز که هدف آن رسیدن به یک نقطه است)، در بسیاری از مسایل دیگر طراح با روش سعی و خطا سیگنال تقویت‌ای عینی‌ای^۲ را که خواست ذهنی‌ی^۳ او را ارضا می‌کند می‌یابد. با این حال معمولاً سیگنال تقویت دست‌یافته با آن سیگنال تقویت بهینه‌ی ناشناخته -که همان‌طور که از

^۱ نتایج این فصل در [Farahmand05a] آمده است.

^۲ Objective

^۳ Subjective

نام‌اش برمی‌آید، طراح از شکل دقیق آن آگاه نیست- تفاوت‌هایی دارد. این تفاوت در سیگنال تقویت باعث تفاوت در تابع ارزش عامل می‌شود و در نتیجه سیاست عامل در مواجهه با محیط تغییر می‌کند. در این فصل فرض می‌کنم که میزان این خطا به صورت کران بالایی مشخص شده است و هدف پیدا کردن کران‌های برای تغییر در تابع ارزش و تغییر در سیاست است.

نتایج این فصل کاربردهای متفاوتی می‌تواند داشته باشد. اینک چند سناریوی ممکن را شرح

می‌دهم:

○ در [Harati02] روش‌هایی برای یادگیری تقویتی سیستم‌های چند-عامله هم‌کار پیش‌نهاد شد. نشان داده شد که با استفاده از معیاری به نام خبرگی می‌توان سرعت یادگیری را بسیاری بالا برد. اما نکته‌ی مهم این بود که در روش‌های تشخیص خبرگی امکان خطا وجود داشت و در نتیجه گاهی سیگنال تقویت به اشتباه به عامل‌های مختلف تخصیص داده می‌شد. میزان این خطا به میزان محافظه‌کار بودن طراحی ما بستگی دارد؛ می‌توان کاملاً محافظه‌کار بود و سیگنال خطای دقیق‌تری داشت ولی در عوض سرعت یادگیری پایین‌تر باشد و یا این که کم‌تر محافظه‌کاری کرد و گاهی سیگنال خطای نادقیق‌ای پیش‌نهاد کرد و سرعت یادگیری را بالا برد ولی در عوض امکان رسیدن به تابع ارزش و در نتیجه سیاست اشتباه به وجود آید. با توجه به میزان اشتباه در سیاست قابل قبول و با استفاده از نتایج این فصل می‌توان میزان محافظه‌کاری روش تعیین خبرگی را مشخص کرد.

○ در بسیاری از مسایل هدف نهایی خوش‌تعریف است - مثلاً رسیدن به نقطه‌ای خاص در ماز، پایداری آونگ معکوس، گل زدن به حریف در فوتبال و گاه به دلیل آن که احتمال رسیدن به حالت هدف بسیار کم است، طراح ممکن است تعدادی هدف میانی - به صورت توابع‌ای جمع شونده به سیگنال تقویت اصلی- برای مساله طراحی کند. این هدف‌های میانی به پاسخ رسیدن را راحت‌تر و سریع‌تر می‌کند. اما در عوض احتمال انتخاب هدف میانی نادرست نیز وجود دارد که ممکن است باعث پایین آمدن کیفیت کار شود. در صورتی که طراح بتواند به نوعی کران خطای احتمالی در

اضافه کردن این هدف‌های میانی را به دست بیاورد می‌تواند از نتایج این فصل برای پیش‌بینی‌ی خطای تابع ارزش و سیاست عامل استفاده کند.

○ نتایج این فصل علاوه بر طراحی سیگنال تقویت، در پیش‌بینی‌ی کیفیت به کارگیری‌ی یک تقریب‌زن تابع مانند شبکه‌ی عصبی در یادگیری تقویتی نیز قابل به کارگیری‌اند. اگر بدانیم که حداکثر خطای تقریب تابع ارزش چه میزان است می‌توان کران‌ای بر خطای سیاست عامل نوشت. چند پژوهش مشابه با این نتایج این فصل وجود دارد. در [Williams94] نتایج‌ای در مورد تاثیر خطای تخمین تابع ارزش بر پاداش بازگشت‌یافته‌ی عامل آمده است. برخلاف پژوهش فعلی آن‌ها بررسی‌ی بر روی تاثیر خطای سیگنال تقویت انجام نداده‌اند و مهم‌تر از آن تاثیر این خطا بر سیاست عامل را -که معیاری رفتارگرایانه از کارکرد عامل است- نیز مطالعه نکرده‌اند ([Sing94] را نیز ببینید). در [Ng99] شرط ناوردایی^۴ سیگنال تقویت اثبات شده است: به ازای چه خانواده توابع سیگنال تقویت، سیاست عامل تغییری نمی‌کند. نتایج کارشان دست‌طراح را در ایجاد سیگنال تقویت برای شکل‌دهی^۵ و سرعت بخشیدن باز می‌کند. اما گاهی دقیق بودن و تغییر نکردن شرط سخت و محدودکننده‌ای است. ممکن است حد اندکی از خطا در سیاست قابل قبول باشد که در چارچوب معرفی شده در آن پژوهش جای نمی‌گیرد.

در این فصل ابتدا به بررسی‌ی تاثیر خطای سیگنال پاداش بر تابع ارزش می‌پردازیم و به صورت ریاضی کران بالایی برای حداکثر تغییر آن تابع به دست می‌آوریم و سپس تاثیر این خطا را بر رفتار عامل به صورت کران‌ای احتمالی به دست می‌آوریم و پس از آن نتایج حاصل را در مساله‌ای کلی نمایش می‌دهیم.

۸-۲) تاثیر عدم قطعیت سیگنال تقویت بر توابع ارزش

در این بخش به بررسی‌ی تاثیر عدم قطعیت بر توابع $Q(s,a)$ و هم‌چنین سیاست $\pi(s)$ ناشی از عدم قطعیت القا شده می‌پردازیم. با دانستن تاثیر آن بر $\pi(s)$ ، می‌توان رفتار سیستم ایده‌آل و سیستم با عدم قطعیت را با هم مقایسه کرد.

^۴ Invariance

^۵ Shaping

فرض کنید که s بیانگر حالت عامل در محیطی مارکوف باشد. عامل عمل a را در محیط انجام می‌دهد و با توجه به تابع انتقال حالت محیط به حالت جدید s' می‌رود و در این گذر تقویت $R(s, a, s')$ را دریافت می‌کند. تابع ارزش بهینه‌ی حالت به صورت زیر تعریف می‌شود:

$$V^*(s) = \max_{\pi} E_{s \sim \pi} \left[\sum_{k=0}^{\infty} \gamma^k r_k \mid s_0 = s \right] \quad (۸-۱)$$

که در آن r_k سیگنال تقویت دریافتی در k امین انتقال حالت سیستم است. امید ریاضی بر روی همه‌ی مسیرهای ممکن‌ای که از $s_0 = s$ آغاز می‌شوند و با سیاست بهینه ادامه می‌یابند گرفته می‌شود. برای سادگی فرمول‌بندی، گاهی از بیان تابع s بودن تابع ارزش صرف‌نظر می‌کنم و آن را بدین‌گونه می‌نویسم:

$$V^* = \max_{\pi} E_{s \sim \pi} \left[\sum_{k=0}^{\infty} \gamma^k r_k \right] \quad (۸-۲)$$

فرض کنید که r_k سیگنال تقویت اصلی ولی ناشناخته‌ای باشد که منجر به پاسخ بهینه‌ی مساله می‌شود (از دیدگاه طراح و نه از دیدگاه عینی که به هر حال فرمول‌بندی یادگیری تقویتی آن به پاسخ بهینه‌ی عینی می‌رسد) و r'_k نیز سیگنال تقویت خطاداری باشد که طراح آن را مشخص کرده است. فرض کنید که r'_k دارای خطایی با نرم محدود باشد، یعنی:

$$r'_k = r_k + \Delta r_k \quad (۸-۳)$$

$$\|\Delta r_k\|_{\infty} < \varepsilon \quad (۸-۴)$$

اندازه‌ی خطا بیانگر عدم قطعیت طراح نسبت به درستی سیگنال طراحی شده‌اش است.

اینک می‌توان تابع ارزش سیگنال تقویت جدید را بدین‌گونه نوشت:

$$V'^*(s) = \max_{\pi} E_s \left[\sum_{k=0}^{\infty} \gamma^k (r_k + \Delta r_k) \right] < \max_{\pi} E_s \left[\sum_{k=0}^{\infty} \gamma^k (r_k + \|\Delta r_k\|_{\infty}) \right] =$$

$$V^* + \varepsilon E_s \left[\sum_{k=0}^{\infty} \gamma^k \right] = V^* + \frac{\varepsilon}{1-\gamma} \quad (۸-۵)$$

که در آن $V'(s)$ تابع ارزش عامل با سیگنال تقویت خطا دار است. به همین صورت می توان نوشت:

$$V'^*(s) = E_s \left[\sum_{k=0}^{\infty} \gamma^k (r_k + \Delta r_k) \right] > E_s \left[\sum_{k=0}^{\infty} \gamma^k (r_k - \|\Delta r_k\|_{\infty}) \right] = V^* - \varepsilon E_s \left[\sum_{k=0}^{\infty} \gamma^k \right] = V^* - \frac{\varepsilon}{1-\gamma}$$

$$(۸-۶)$$

و در نتیجه:

$$V^* - \frac{\varepsilon}{1-\gamma} < V' < V^* + \frac{\varepsilon}{1-\gamma} \quad (۸-۷)$$

و یا

$$\|\Delta V\|_{\infty} < \frac{\varepsilon}{1-\gamma}$$

که در آن $\Delta V = V'^* - V^*$.

به همین صورت در مورد $Q(s, a)$ نیز می توان عبارتهای مشابهی به دست آورد. فرض کنید

که عامل در حالت s باشد و با سیاست $\pi'(s)$ - که تابعی از $Q'(s, a)$ است - عملی را انتخاب کند.

محیط با احتمال $P_{sa} : (s, a) \mapsto s'$ به حالت جدید s' می رود. پس خواهیم داشت:

$$Q'^*(s, a) = \max_{\pi} E_{s' \sim P_{sa}^{\pi}} [r'_k + \gamma V'^*(s')] \quad (۸-۹)$$

حال به مانند قبل می توان نوشت:

$$\begin{aligned} Q'^*(s, a) &= \max_{\pi} E_{s' \sim P_{sa}^{\pi}(s)} [r'_k + \gamma V'^*(s')] = E_{s'} [r(s) + \Delta r(s) + \gamma V'^*(s')] \\ &< E_{s'} \left[r(s) + \Delta r(s) + \gamma \left(V^*(s') + \frac{\varepsilon}{1-\gamma} \right) \right] = E_{s'} \left[\left(r(s) + \gamma V^*(s') \right) + \left(\Delta r(s) + \frac{\varepsilon}{1-\gamma} \right) \right] \end{aligned} \quad (۸-۱۰)$$

می بینید که کران بالا به دو قسمت تقسیم شده است: قسمت اول، همان $Q^*(s, a)$ است و قسمت دوم، ناشی از عدم قطعیت - یعنی:

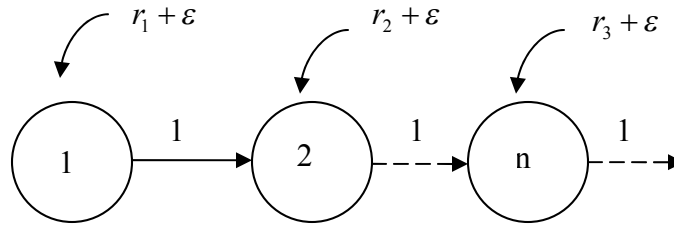
$$\begin{aligned} Q'^*(s, a) &< Q^*(s, a) + E_{s'} \left[\Delta r(s) + \frac{\gamma \varepsilon}{1-\gamma} \right] = Q^*(s, a) + \frac{\varepsilon}{1-\gamma} + E_{s'} [\Delta r(s)] \\ &< Q^*(s, a) + \frac{\gamma \varepsilon}{1-\gamma} + \varepsilon = Q^*(s, a) + \frac{\varepsilon}{1-\gamma} \end{aligned} \quad (۸-۱۱)$$

و به همین ترتیب می توان نشان داد که:

$$\begin{aligned} Q'^*(s, a) &= E_{s'} [r(s) + \Delta r(s) + \gamma V'^*(s')] > E_{s'} \left[r(s) + \Delta r(s) + \gamma \left(V^*(s') - \frac{\varepsilon}{1-\gamma} \right) \right] \\ &= E_{s'} \left[\left(r(s) + \gamma V^*(s') \right) + \left(\Delta r(s) - \frac{\gamma \varepsilon}{1-\gamma} \right) \right] > Q^*(s, a) - \frac{\gamma \varepsilon}{1-\gamma} + E_{s'} [\Delta r(s)] \\ &> Q^*(s, a) - \frac{\gamma \varepsilon}{1-\gamma} - \varepsilon = Q^*(s, a) - \frac{\varepsilon}{1-\gamma} \end{aligned} \quad (۸-۱۲)$$

و در نتیجه داریم:

$$\|\Delta Q(s, a)\|_{\infty} < \frac{\varepsilon}{1-\gamma} \quad (۸-۱۳)$$



شکل ۸-۱. نمونه‌ای از مسأله‌ی MDP که کران‌های ارایه‌شده دقیق‌اند.

که مشابه با قبل $\Delta Q = Q^* - Q$.

ممکن است به نظر آید که کران‌های به دست آمده بسیار محافظه‌کارانه هستند. با این‌وجود مسائلی وجود دارد که این کران‌ها دقیقاً تخمین درست خطای تابع ارزش هستند. مثلاً در شکل ۸-۱ که در آن انتقال حالت برای حرکت چپ به راست ۱ است و سیگنال تقویت همیشه با خطای ثابت‌ای معادل با ϵ آلوده شده است، کران به دست آمده معادل میزان دقیق خطا خواهد بود.

۸-۳) تاثیر عدم قطعیت سیگنال تقویت بر سیاست عامل

در این بخش تاثیر تغییر توابع ارزش بر سیاست عامل را بررسی می‌کنیم. اگر عامل به سیگنال تقویت بهینه‌ی ناشناخته دسترسی داشته باشد، سیاست بهینه‌ی عامل $\pi(s, a)$ خواهد بود. حال سوال این است که به هنگام دریافت سیگنال تقویت خطادار ارتباط سیاست جدیدش $\pi'(s, a)$ با سیاست بهینه‌ی ناشناخته چگونه خواهد بود. با این کار می‌توانیم بفهمیم که از دیدگاه رفتارگرا تاثیر عدم قطعیت در سیگنال تقویت چگونه ظاهر می‌شود. برای مقایسه این دو چگالی احتمال از احتمال نسبی انتخاب عمل در حالت‌ای یک‌سان استفاده می‌کنیم: $\pi'(s, a) / \pi(s, a)$. اگر این مقدار به یک نزدیک بود، آن‌گاه می‌توان گفت که رفتار عامل خطادار چندان تفاوتی با عامل بهینه‌ی اصلی ندارد ولی اگر این مقدار تفاوت زیادی داشت (خیلی کوچک یا خیلی بزرگ بود) آن‌گاه رفتار عامل مطمئناً تغییر کرده است. در نتایج پیش رو فرض شده است که عامل از روش انتخاب عمل بولتزمن که از همه‌ی اطلاعات تابع ارزش حالت-عمل استفاده می‌کند برای انتخاب عمل بهره می‌گیرد. در روش انتخاب بولتزمن داریم:

$$P\left\{\text{selected action} = a^*(s) \middle| Q(s, a), \text{Boltzman action selection, } T\right\} = \frac{\exp\left(\frac{Q(s, a^*)}{T}\right)}{\sum_a \exp\left(\frac{Q(s, a)}{T}\right)} = \frac{N}{D} \quad (۸-۱۴)$$

پس برای عامل خطادار خواهیم داشت^۶:

$$P\left\{a^*(s) \middle| Q'(s, a)\right\} = \frac{\exp\left(\frac{Q'(s, a^*)}{T}\right)}{\sum_a \exp\left(\frac{Q'(s, a)}{T}\right)} = \frac{\exp\left(\frac{Q(s, a^*)}{T}\right) \exp\left(\frac{\Delta Q(s, a^*)}{T}\right)}{\sum_a \left[\exp\left(\frac{Q(s, a)}{T}\right) \cdot \exp\left(\frac{\Delta Q(s, a)}{T}\right)\right]} = \frac{N'}{D'} \quad (۸-۱۵)$$

برای محاسبه‌ی کران بالا و پایینی برای این احتمال، می‌توان کران بالا و پایین صورت و مخرج را محاسبه کرد و سپس با استفاده از آن‌ها، کران کل را به دست آورد.

$$N' = \exp\left(\frac{Q(s, a^*)}{T}\right) \exp\left(\frac{\Delta Q(s, a^*)}{T}\right) = N \cdot \exp\left(\frac{\Delta Q(s, a^*)}{T}\right) \quad (۸-۱۶)$$

اگر تعریف کنیم:

$$\rho = \exp\left(\frac{\|\Delta Q(s, a)\|_{\infty}}{T}\right) \quad (۸-۱۷)$$

آن‌گاه

$$\rho^{-1} N \leq N' \leq \rho N \quad (۸-۱۸)$$

^۶ توجه شود که شروط بدیهی‌ای مانند دما و انتخاب نوع عمل بولتزمن -به جز حالتی که یادآوری آن اهمیت داشته باشد- در فرمول‌های بعدی حذف شده است.

به همین صورت برای D' نیز می توان نوشت:

$$\sum_a \exp\left(\frac{Q(s,a)}{T}\right) \cdot \rho^{-1} \leq D' = \sum_a \exp\left(\frac{Q(s,a)}{T}\right) \cdot \exp\left(\frac{\Delta Q(s,a)}{T}\right) \leq \sum_a \exp\left(\frac{Q(s,a)}{T}\right) \cdot \rho$$

(۸-۱۹)

و در نتیجه:

$$\rho^{-1} D \leq D' \leq \rho D$$

(۸-۲۰)

اکنون می توان نوشت:

$$\frac{\min(N')}{\max(D')} \leq P\{a^*(s)|Q'(s,a)\} \leq \frac{\max(N')}{\min(D')}$$

(۸-۲۱)

پس:

$$\rho^{-2} P\{a^*(s)|Q(s,a)\} \leq P\{a^*(s)|Q'(s,a)\} \leq \rho^2 P\{a^*(s)|Q(s,a)\}$$

(۸-۲۲)

و یا:

$$\rho^{-2} \leq \frac{P\{a^*(s)|Q'(s,a)\}}{P\{a^*(s)|Q(s,a)\}} \leq \rho^2$$

(۸-۲۳)

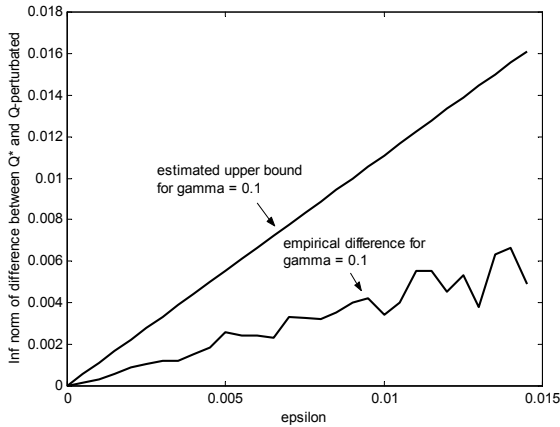
حال با در نظر گرفتن کران بخش پیش داریم:

$$\rho = \exp\left(\frac{\varepsilon}{(1-\gamma)T}\right) \quad (۸-۲۴)$$

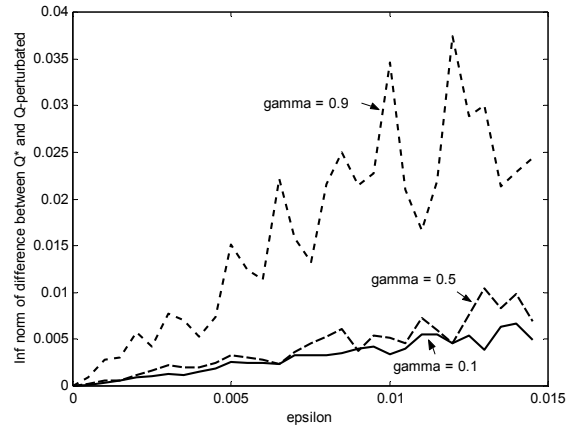
۸-۴) آزمایش‌ها

در این بخش درستی کران‌های محاسبه شده را با مقایسه با کران واقعی‌ای که از طریق شبیه‌سازی به دست آمده است نمایش می‌دهیم. برای این کار به ازای مقادیر مختلف خطای ε (در ۴-۸) میزان خطای ایجاد شده را $\|\Delta Q\|_\infty$ را با خطای پیش‌بینی در کران‌های به دست آمده مقایسه می‌کنیم. مشابه با همین کار را برای تغییر در سیاست‌ها نیز انجام می‌دهیم. با این کار گستره‌ی قابل استفاده کران‌های به دست آمده نیز تا حدودی مشخص می‌شود.

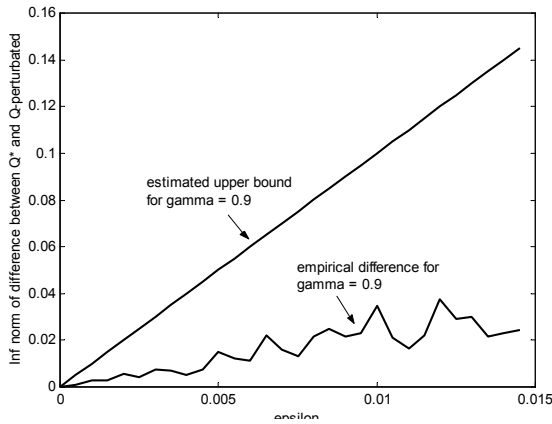
برای مقایسه می‌توان یکی از مسایل معمول و متداول‌ای را که با استفاده از روش یادگیری تقویتی به حل آن مبادرت می‌کنند انتخاب کرد. به عنوان مثال می‌توان به مساله‌ی مسیریابی در ماز، متعادل‌سازی ی‌آونگ معکوس، و یا حتی مسایل پیچیده‌تری مانند یادگیری رفتارها در شبیه‌سازی بازی فوتبال اشاره کرد. با این وجود، به دلیل آن‌که خصوصیات ذاتی مساله تأثیری بر نتایج نداشته باشد و همچنین به دلیل آن‌که پیچیدگی‌هایی چون انتخاب حالت و نوع بازنمایی - که در مساله‌ای چون یادگیری رفتار در فوتبال مطرح است - باعث سخت شدن مساله حتی برای نسخه‌ی بدون خطای آن نشود از مسایل متداول استفاده نشده است، بلکه از مساله‌ای فرضی - ولی کلی - که به طور اتفاقی طراحی شده بهره گرفته شده است. بدین صورت که محیط‌ای در نظر گرفته شده که دارای احتمال انتقال حالت $(P_{sa} : (s, a) \mapsto s')$ تصادفی است و پاداش انتقال به هر حالت جدید $(R(s))$ نیز به صورت تصادفی مشخص شده. در واقع این مساله معادل حرکت بهینه بر روی گراف‌ای تصادفی با یال‌های وزن دار است که وزن یال‌ها متناسب با زمان عبور از آن‌ها مشخص می‌شود و می‌توان آن را به عنوان کلی‌ترین حالت مساله‌ی یادگیری تقویتی (و در نتیجه برنامه‌ریزی پویا) در نظر گرفت. لازم است تأکید کنم که این مساله درست به همان اندازه‌ای که تئوری‌های این فصل احتیاج دارند کلی است و استفاده از مسایل دیگر یا زیر پا گذاشتن فرض‌هایی از تئوری است (مثلاً با POMDP شدن مساله در حالی که فرمول‌بندی احتیاج به MDP بودن آن دارد) و یا این‌که قیدهایی بی‌فایده - جز فایده‌ی نمایش‌ای - به مساله اضافه می‌کند.



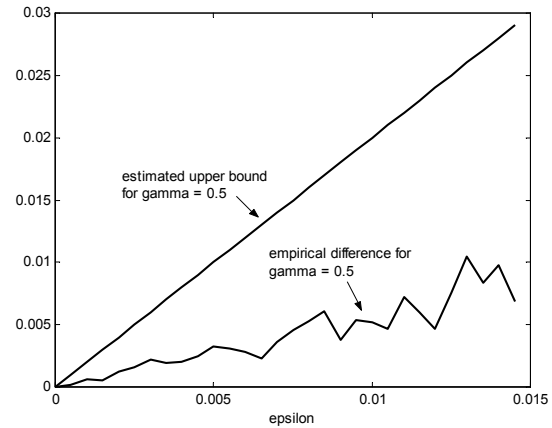
شکل ۸-۳. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.1$



شکل ۸-۲. خطای $\| \Delta Q(s, a) \|_{\infty}$ به ازای مقادیر γ مختلف

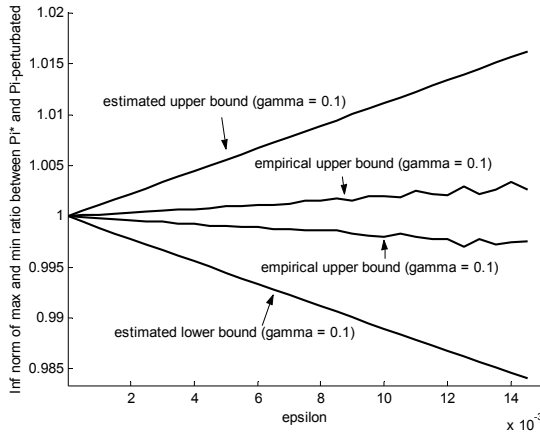


شکل ۸-۵. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.9$

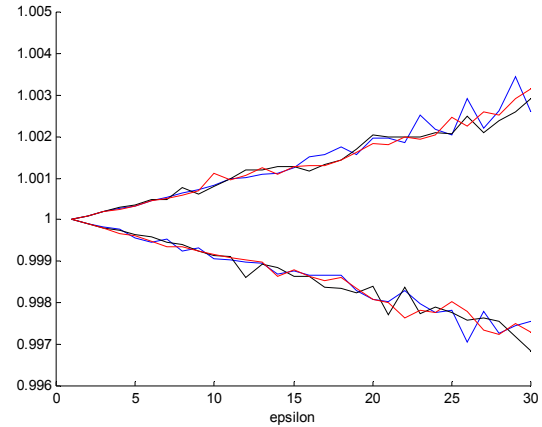


شکل ۸-۴. مقایسه بین خطای $\| \Delta Q(s, a) \|_{\infty}$ مشاهده شده و کران به دست آمده به ازای $\gamma=0.5$

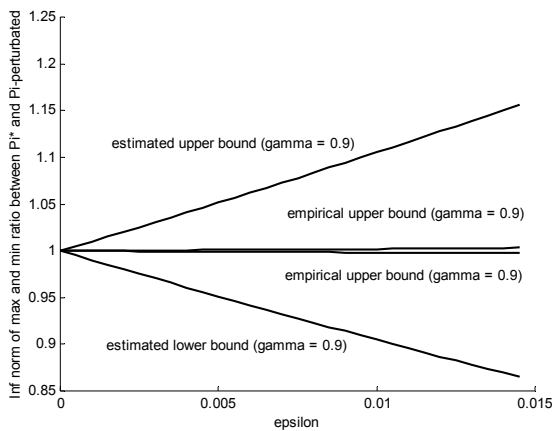
برای نتایجی که در این مقاله گزارش می‌شود، تعداد حالت‌ها (راس‌های گراف) را برابر ۳۰ در نظر می‌گیرم، تعداد عمل‌های ممکن در هر حالت را نیز ۵ قرار می‌دهم و احتمال انتقال را نیز از توزیعی یک‌نواخت انتخاب می‌کنیم با این شرط که در نهایت تابع انتقال احتمالی صحیحی‌ای تولید شود (یعنی مجموعه احتمال خروج از هر حالت-عمل برابر ۱ باشد). تابع تقویت را نیز از توزیع یک‌نواختی در بازه‌ی $(-1, 1)$ انتخاب کرده‌ایم. هم‌چنین برای محاسبه‌ی سیاست در هر حالت نیز از انتخاب بولترمان با دمای $T = 1$ استفاده شده است. قابل ذکر است که برای بالا بردن دقت، به ازای هر مقدار ϵ ، این آزمون ۱۰ مرتبه انجام شده است و بالاترین کران به دست آمده در نمودارها آمده است.



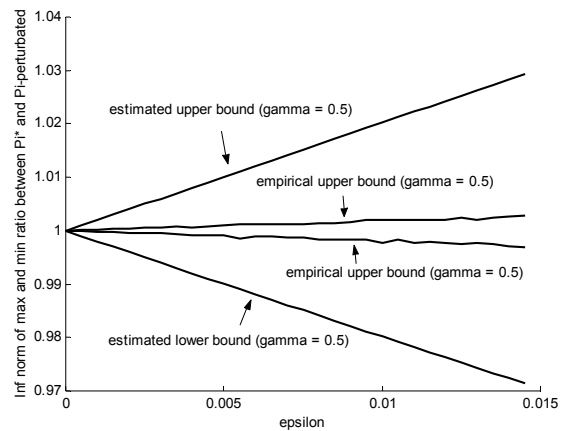
شکل ۷-۸. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.1$



شکل ۷-۹. کران بالا و پایین نسبت احتمالات عامل با سیگنال تقویت نادقیق به احتمالات عامل با سیگنال تقویت اصلی به ازای مقادیر مختلف γ (آبی: $\gamma=0.1$ ، مشکی: $\gamma=0.5$ ، قرمز: $\gamma=0.9$).



شکل ۸-۸. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.9$



شکل ۸-۹. مقایسه بین نسبت احتمالات مشاهده شده و کران‌های به دست آمده به ازای $\gamma=0.5$

در شکل ۲-۸ میزان خطای عملی بین تابع ارزش دارای خطای سیگنال تقویت و تابع ارزش اصلی به ازای مقادیر مختلف ϵ و γ آمده است. همان‌طور که انتظار می‌رفت هر چقدر میزان خطای تابع تقویت بیش‌تر باشد، میزان خطای تابع ارزش نیز بیش‌تر خواهد بود. علاوه بر آن هنگامی که γ بزرگ‌تر باشد تاثیر خطا نیز بیش‌تر است. در شکل‌های ۳-۸ تا ۵-۸ کران خطایی که در عمل رخ داده با کران‌ای که در این مقاله به دست آوردیم (رابطه ۱۳-۸) مقایسه شده است. مشاهده می‌شود که در صورتی که γ و ϵ کوچک باشد، دقت کران بیش‌ترست. در شکل ۶-۸ کران‌های بالا و پایین نسبت احتمال انتخاب

عمل در حالت دارای سیگنال تقویت مخدوش شده با حالت اصلی در مثال مورد بررسی مقایسه شده است (به رابطه ۸-۲۳ نگاه کنید). نکته‌ی جالب این است که این مقادیر چندان به γ وابسته نیستند. به مانند قبل، این کران‌های عملی با آنچه محاسبات‌مان مشخص کرده است مقایسه می‌کنیم. این مقایسه به ازای مقادیر مختلف γ در شکل‌های ۷-۸ تا ۸-۹ آمده است. مشاهده می‌شود که به ازای مقادیر γ و یا ε بزرگ کران‌های‌مان بسیار محافظه‌کارانه هستند – گرچه به یاد بیاورید که کران‌های‌ی به دست آمده برای بعضی از مسایل کران دقیق هستند.

۸-۵) نتیجه‌گیری

در این فصل تاثیر خطای سیگنال تقویت بر ارزش هر حالت و همچنین احتمال انتخاب عمل‌های هر حالت به صورت کران‌هایی به دست آمد. بدین صورت می‌توان تخمین‌ای کلی از این که در صورت خطا و عدم قطعیت در سیگنال تقویت به چه میزان عامل‌مان دچار تغییر خواهد شد به دست آورد. اما همان‌طور که در شبیه‌سازی‌ها دیده شد، این کران‌ها – به دلیل طبیعت کلی‌شان – محافظه‌کارانه هستند و به ازای مقادیر بزرگ خطا عملاً بی‌فایده خواهند بود. به همین دلیل عقلایی می‌نماید که بتوان با توجه به ساختار خود مساله کران‌های کم محافظه‌کارانه‌تری به دست آورد. از طرف دیگر ممکن است بتوان با تحلیل‌های دقیق‌تر ریاضی و با استفاده از کران‌های متداول در احتمالات (مانند کران Hoeffding و ...) به کران‌های دقیق‌تری – که در اکثر موارد و نه همه‌ی موارد درست‌اند^۷ – رسید.

^۷ مشابه با ایده‌ی Probably Approximately Correct (PAC) که در تئوری‌های یادگیری ماشینی آماری به کار می‌روند [Vidyasagar03].

۹) نتیجه‌گیری و پیش‌نهاد

در این پایان‌نامه چندین روش نوین برای خودکارسازی طراحی عامل‌های هوشمند رفتارگرا ارائه دادم. برخلاف بسیاری از پژوهش‌های پیشین، سعی در ارائه‌ی نگاه‌ای چندجانبه و دقیق به مساله‌ی طراحی عامل رفتارگرا کردم. تا حد ممکن مساله را به قالب ریاضی در آوردم و آن را به طور دقیق تحلیل کردم. این موضوع به خصوص در مورد یادگیری ساختار و یادگیری رفتار نمود بیش‌تری داشت. نتایج تحلیل‌های انجام شده منجر به چندین روش‌های مختلف یادگیری و تکامل شد. این روش‌ها به طور فهرست‌وار بدین شرح‌اند:

- یادگیری ساختار
 - بازنمایی مرتبه صفر
 - بازنمایی مرتبه یک
- یادگیری رفتار
- تکامل رفتار
- اشتراک سازگاری یک‌نواخت
- اشتراک سازگاری ارزش‌محور
- الگوریتم ممیتیک

در فصل مربوط به یادگیری ساختار (فصل ۳) دو نوع بازنمایی دانش مختلف پیش‌نهاد شد که با استفاده از آن، عامل می‌تواند اطلاعات دریافتی از محیط را ثبت کرده و به‌ترین ساختار ممکن را بیابد. نشان دادم که چگونه می‌توان ارزش عامل را به اجزای ساده‌تری تجزیه کرد و متناسب با آن تجزیه، روش‌ای برای تخصیص اعتبار به رفتارها و چگونگی به روز آوری آن‌ها را معرفی کردم.

در فصل مربوط به یادگیری رفتار (فصل ۴) با معرفی عمل مجازی‌ای به نام NA به رفتارها اجازه دادم تا تخمین‌ای از عمل‌کرد رفتارهای دیگر به دست آورند و متناسب با آن تخمین، تصمیم بگیرند که آیا به‌تر است خودشان کنترل عامل را در اختیار بگیرند یا اجازه‌ی چنین کاری را به رفتارهای دیگر بدهند. با اضافه کردن این عمل مجازی، امکان هم‌کاری بین عامل‌ها (رفتارها) بیش‌تر می‌شود. روی‌کرد دیگری که در این پایان‌نامه مطرح شد، استفاده از تکامل هم‌کارانه‌ی رفتارها بود. روش‌ای برای تکامل رفتارها ارایه دادم (فصل ۵) که ماجولارتر و قابل استفاده‌ی مجددتر از روش‌های یک‌پارچه‌ی پیشین است. دو روش مختلف اشتراک سازگاری بین جمعیت‌های تکامل‌یابنده معرفی کرد که هر کدام برگرفته از فلسفه‌ای خاص بودند.

در نهایت تفسیری جدید و متفاوت از مم ارایه دادم (فصل ۵) که اجازه می‌دهد به جای هدر رفتن دانش کسب‌شده در طول زندگی عامل، آن دانش به صورت سنت‌های فرهنگی قابل استفاده‌ی مجدد برای عامل‌های جدید در می‌آید. نتایج حاصل در آزمایش‌های انجام گرفته موفقیت‌آمیز بود.

با در اختیار داشتن این مجموعه از روش‌ها، طراح با دست باز می‌تواند یک یا چند نوع از این روش‌ها را با توجه به دانش‌اش نسبت به مساله اختیار کند. به عنوان نمونه اگر او جعبه‌ابزاری از رفتارهای مناسب در اختیار داشته باشد می‌تواند از یکی از دو روش یادگیری ساختار بهره بگیرد و اگر نه، از یادگیری یا تکامل برای تولید رفتارهای جدید استفاده کند. به نظر می‌رسد روش‌های پیش‌نهادی تا حد قابل قبولی روند طراحی‌ی خودکار عامل هوشمند را ساده کنند. نتایج به‌کارگیری این روش‌ها در فصل آزمایش‌ها (فصل ۶) بر روی دو مساله‌ی مختلف به نمایش گذاشته شده است و منجر به نتایج قابل رقابت و حتی به‌تر از معماری‌های طراحی شده توسط انسان بود.

علاوه بر تحلیل‌هایی که به طور مستقیم به روش‌ای برای طراحی منجر شدند، تحلیل‌های دیگری نیز صورت گرفت که درک ما را از مساله‌ی یادگیری در ساختاری سلسله‌مراتبی بالا می‌برد. یکی از این تحلیل‌ها، بررسی‌ی احتمال کنترل‌کننده بودن و به دست آوردن توزیع احتمال تعداد دفعه‌ی کنترل‌کننده بودن برای معماری‌های رقابتی‌ی موازی‌ای مانند PPSSA است (فصل ۷). نتایج حاصل با این‌که شاید به طور مستقیم در روند طراحی به کار نروند، اما درک طراح از مساله را افزایش می‌دهد.

دسته‌ی دیگر تحلیل‌ها، بررسی‌ی اثر عدم قطعیت سیگنال تقویت بر تابع ارزش و تابع سیاست عامل بود (فصل ۸). با محاسبات ساده‌ای نشان دادم که عدم قطعیت چگونه می‌تواند رفتار عامل را عوض کند. این موضوع هنگامی اهمیت پیدا می‌کند که بدانیم در بسیاری از مسایل دنیای واقعی، دانش ما نسبت به سیگنال تقویت دقیق نیست و آن سیگنال ممکن است دارای خطا باشد. با این‌که این پژوهش تا حد ممکن وسیع و گسترده بوده است، اما همچنان مسایل مهم‌ای برای بررسی باقی مانده. به عنوان مثال می‌توان به موارد زیر اشاره کرد:

○ **تعمیم ایده‌های مطرح‌شده به ساختارهای پیچیده‌تر.**

در این پایان‌نامه فرض کرده‌ام که معماری رفتارگرایی سلسله‌مراتبی مان دارای ساختاری اکیدا موازی است. با این‌که بسیاری از رفتارها را می‌توان با همین معماری ایجاد کرد، اما به نظر می‌رسد گسترش این ایده‌ها به معماری‌های کلی‌تر مفید باشد. گمان می‌کنم در صورت‌ای که بتوان ارزش معماری را به صورت ارزش اجزای ساده‌تری در آورد، آنگاه یادگیری ساختار آن معماری‌ها نیز ممکن باشد.

○ **استخراج خودکار فضای حالت و عمل رفتارها**

در این پایان‌نامه فرض کرده بودم که بُعد فضای حالت و عمل رفتارها از پیش دانسته است. خودکاری‌سازی این فرآیند و استخراج خودکار ویژگی‌هایی که می‌توانند برای تصمیم‌گیری مطلوب باشند به نظر مساله‌ی دشواری می‌نماید. پژوهش‌های اولیه‌ای (و البته نه قطعی) که در این زمینه انجام داده‌ام نشان می‌دهد که استفاده از روش‌های خوشه‌بندی^{۱۸۰} متداول نمی‌تواند خیلی موثر باشد و لازم است الگوریتم‌های خوشه‌بندی جدیدی پیشنهاد شوند که علاوه بر در نظر گرفتن ویژگی‌های توپولوژیک حالت‌های تجربه شده، از ارزش آن حالت‌ها نیز بهره بگیرد. دو پژوهش‌ای که تا حدی به این موضوع پرداخته‌اند [Hengst02] و روش ادغام حالت [Engel01] است.

¹⁸⁰ Clustering

○ هم‌گرایی روش‌های یادگیری

در این پایان‌نامه اثبات‌ای از هم‌گرایی توابع ارزش به مقادیر بهینه‌ی عمومی و یا حتی محلی‌ی خود انجام نشد. پژوهش راجع به این‌که آیا روش‌های یادگیری پیشنهادشده به بهینه‌ای (عمومی یا محلی) هم‌گرا می‌شوند یا خیر جالب توجه خواهد بود.

○ طراحی سیگنال تقویت

مسالهِی طراحی سیگنال تقویت، مسالهِی اساسی در یادگیری تقویتی است و منحصر به این پایان‌نامه نمی‌شود. این‌که آیا روش‌ای کلی و سیستماتیک برای طراحی سیگنال تقویت وجود دارد یا خیر مشخص نیست. تنها چند پژوهش محدود درباره‌ی این موضوع انجام شده است که می‌توان به [Dorigo94]، [Dorigo97]، [Ng99] و [Farahmand05a] اشاره کرد. یکی از ایده‌های مطرح شده، ایجاد خودکار سیگنال تقویت با استفاده از شمایی کلی از آن‌چه عامل می‌بایست رفتار کند است. نتایج مقدماتی‌ای در این باره به دست آمده که البته در این پایان‌نامه گزارش نشده است.

○ تجرید دانش سطح‌پایین کسب شده و استفاده از آن در تصمیم‌گیری‌های سطح‌بالا تر

یکی از زمینه‌های جالب توجه برای ادامه‌ی این پژوهش، تجرید دانش کسب شده توسط رفتارها و استفاده از آن در لایه‌های سطح‌بالا تر تصمیم‌گیری است. تلاش‌های اولیه‌ای در این زمینه انجام شده است که نتایج جالب توجه‌ای داشته است. این نتایج تاکنون گزارش نشده‌اند.

(۱۰) مراجع

1. [Arkin98] R. C. Arkin, Behavior-Based Robotics, *MIT Press*, CA, 1998.
2. [Boyan95] J. Boyan and A. W. Moore, "Generalization in reinforcement learning: safely approximating the value function," in G. Tesauro, D. S. Touretzky, and T. K. Leen (Eds.), *Advances in Neural Information Processing System 7*, MIT Press, Cambridge MA, 1995.
3. [Brooks86] R. A. Brooks, "A robust layered control system for a mobile robot," *IEEE Journal of Robotics and Automation R.A-2*, 1986, pp. 14-23.
4. [Brooks89] R. A. Brooks, "A robot that walks: emergent behavior from a carefully evolved network," *Neural Computation* 1(2), 1989, pp. 252-262.
5. [Brooks91a] R. A. Brooks, "New approaches to robotics," *Science* 253, 1991, pp. 1227-1232.
6. [Brooks91b] R. A. Brooks, "Intelligence without reason," in *Proc. Int. Joint Conf. AI*, 1991, pp. 569-595.
7. [Brooks91c] R. A. Brooks, "Intelligence without representation," *Artificial Intelligence*, 47, 1991, pp. 139-159.
8. [Buriol04] L. Buriol, P. M. Franca, P. Moscato, "A new memetic algorithm for the asymmetric traveling salesman problem," *Journal of Heuristics*, 10, pp. 483-506, 2004.
9. [Chernova04] S. Chernova and M. Veloso, "An Evolutionary Approach to Gait Learning for Four-Legged Robots," *IEEE/RSJ Int. Conf. Intelligent Robots and Systems*, 2004.
10. [Cliff93] D.T. Cliff, I. Harvey, and P. Husbands, "Explorations in Evolutionary Robotics," *Adaptive Behavior*, 2, pp. 73-110, 1993.
11. [Dawkins76] R. Dawkins, *The Selfish Gene*, *Oxford University Press*, 1976.
12. [Dayan93] P. Dayan and G. Hinton, "Feudal reinforcement learning," *Advances in Neural Information Processing Systems (NIPS)*, 5, Morgan Kaufmann, San Francisco, CA, 1993, pp. 271-278.
13. [Dorigo94] M. Dorigo and M. Colombetti, "Robot shaping: developing autonomous agents through learning," *Artificial Intelligence*, 71(4), 1994, pp. 321-370.
14. [Dorigo97] M. Dorigo and M. Colombetti, *Robot Shaping: An Experiment in Behavior Engineering*, *MIT Press*, CA, 1997.
15. [Engel01] Y. Engel and S. Mannor, "On finding good state aggregation functions," *Workshop on Hierarchy and Memory in Reinforcement Learning, International Conference on Machine Learning (ICML)*, 2001.
16. [Farahmand04] A. M. Farahmand, M. Nili Ahmadabadi, and B. N. Araabi, "Behavior hierarchy learning in a behavior-based system using reinforcement

- learning,” *IEEE/RSJ Int. Conf. Intelligent Robots and Systems*, Sendai, Japan, 2004, pp. 2050-2055.
17. [Farahmand05a] A. M. Farahmand and M. Nili Ahmadabadi, “The effect of reinforcement signal in reinforcement learning,” *The Computer Society of Iran Computer Conference (CSICC)*, 2005 (in Persian).
 18. [Farahmand05b] A.M. Farahmand, M. Nili Ahmadabadi, B. N. Araabi, “Behavior and hierarchy development in behavior-based systems using reinforcement learning,” *Submitted to Artificial Intelligence*, 2005.
 19. [Federici03] D. Federici, “Combining genes and memes to speed up evolution,” in *Proceedings of International Congress on Evolutionary Computation (CEC)*, 2003.
 20. [Floreano94] D. Floreano, and F. Mondada, “Automatic creation of an agent: genetic evolution of a neural network driven robot,” In D. Cliff, P. Husbands, J.-A. Meyer, and S. Wilson (Eds.), *From Animals to Animats III*, Cambridge, MA: MIT Press, 1994.
 21. [Floreano96] D. Floreano and F. Mondada, “Evolution of homing navigation in a real mobile robot,” *IEEE Transactions on Systems, Man, and Cybernetics - Part B*, 26, pp. 396-407, 1996.
 22. [Floreano99] D. Floreano, S. Nolfi, and F. Mondada, “Co-evolution and ontogenetic change in competing robots,” *Robotics and Autonomous Systems*, 1999.
 23. [Floreano00] D. Floreano and J. Urzelai, “Evolutionary robotics: the next generation,” in T. Gomi (ed.), *Proceedings of Evolutionary Robotics*, III, pp. 231-266, Ontario, 2000.
 24. [Goldberg89] David E. Goldberg, *Genetic Algorithms in Search, Optimization and Machine Learning*, Addison-Wesley Publishing, 1989.
 25. [Gomez97] F. Gomez and R. Miikkulainen, “Incremental evolution of complex general behavior,” *Adaptive Behavior*, 5, pp. 317-342, 1997.
 26. [Harati04] A. Harati and M. Nili Ahmadabadi, “Experimental analysis for knowledge based multiagent credit assignment,” *Neural Information Processing: Research and Development, Research and Development*, Rajapakse, Jagath C.; Wang, Lipo (Eds.), Springer-Verlag, May 2004.
 27. [Harvey93] I. Harvey, P. Husbands, and D. Cliff, “Issues in evolutionary robotics,” in *From Animals to Animats II: Proceedings of the Second International Conference on Simulation of Adaptive Behavior*, J. Meyer, H. L. Roitblat, and S. W. Wilson, Eds. MIT Press-Bradford Books, Cambridge, MA, 1993.
 28. [Hengst02] B. Hengst, “Discovering hierarchies in reinforcement learning with HEXQ,” in Claude Sammut and Achim Hoffmann, (Eds.), *Proceeding of Nineteenth International Conference on Machine Learning*, Morgan-Kaufman, 2002, pp. 243-250.
 29. [Hergenhahn01] B. R. Hergenhahn and M. H. Olson, *An Introduction to Theories of Learning*, Prentice-Hall, 2001.
 30. [Jaakkola94] T. Jaakkola, M. I. Jordan, and S. Singh, “On the convergence of stochastic iterative dynamic programming algorithms,” *Neural Computation*, 6(6), 1994, pp. 1185-1201.

31. [Kaelbling96] L. P. Kaelbling and M. L. Littman, and A. W. Moore, "Reinforcement learning: a survey," *Journal of Artificial Intelligence Research*, 4, 1996, pp. 237-285.
32. [Kohl04] N. Kohl and P. Stone, "Policy gradient reinforcement learning for fast quadrupedal locomotion," *Proceeding of the IEEE International Conference on Robotics and Automation (ICRA)*, 2004.
33. [Koza94] John Koza, "Evolution of a subsumption architecture that performs a wall following task for an autonomous mobile robot via genetic programming," In S.J. Hanson, T. Petsche, M. Kearns, and R.L. Rivest (editors), *Computational Learning Theory and Natural Learning Systems*, The MIT Press, vol. 2, pp. 321-346, 1994.
34. [Krasnogor04] N. Krasnogor and S. Gustafson, "A study on the use of 'Self-Generation' in memetic algorithms," *Natural Computing* 3 (1), pp. 53-76, 2004.
35. [Maes90] P. Maes and R. A. Brooks, "Learning to coordinate behaviors," in: *Proc. AAAI*, Boston, MA, 1990, pp. 796-802.
36. [Mahadevan92] S. Mahadevan and J. Connell, "Automatic programming of behavior-based robots using reinforcement learning," *Artificial Intelligence*, 55, 1992, pp. 311-365.
37. [Matarić92] M. J. Matarić, "Integration of representation into goal-driven behavior-based robots," *IEEE Transactions on Robotics and Automation*, 8(3), 1992, pp. 46-54.
38. [Matarić94] M. J. Matarić, "Reward function for accelerated learning," in: W. W. Cohen and H. Hirsh, (Eds.), *Proc. 8th International Conference on Machine Learning*, Morgan Kaufmann, 1994, pp. 181-189.
39. [Matarić97] M. J. Matarić, "Reinforcement learning in the multi-robot domains," *Autonomous Robots*, vol. 4, no. 1, 1997, pp. 73-88.
40. [Matarić98] M. J. Matarić, "Behavior-based robotics as a tool for synthesis of artificial behavior and analysis of natural behavior," *Trends in Cognitive Science*, vol. 2, no. 3, 1998, pp. 82-87.
41. [Matarić01] M. J. Matarić, "Learning in behavior-based multi-robot systems: policies, models, and other agents," *Cognitive System Research*, special issue on Multi-disciplinary studies of multi-agent learning, Ron Sun, (Ed.), 2(1), 2001, pp. 81-93.
42. [McCallum95a] R. A. McCallum, "Instance-based state identification for reinforcement learning," *Advances in Neural Information Processing Systems (NIPS)*, 1995.
43. [McCallum95b] R. A. McCallum, "Instance-based Utile Distinctions for reinforcement learning with hidden state," *Proceedings of the 12th International Machine Learning Conference (ML)*, 1995.
44. [McGovern01] A. McGovern and A. Barto, "Automatic discovery of subgoals in reinforcement learning using diverse density," in: C. Brodley and A. Danyluk, (Eds.), *Proceeding of the 18th International Conference on Machine Learning*, San Francisco, CA., Morgan Kaufmann, 2001, pp. 361-368.

45. [Merz99] P. Merz and B. Freisleben, "A comparison of memetic algorithms, tabu search, and ant colonies for the quadratic assignment problem," In *Proceedings of the 1999 International Congress of Evolutionary Computation*, pp. 2063-2070, 1999.
46. [Merz00] P. Merz and B. Freisleben, "Fitness landscape analysis and memetic algorithms for the quadratic assignment problem," *IEEE Transactions on Evolutionary Computation*, vol. 4, no. 4, pp. 337-352, 2000.
47. [Michaud98] F. Michaud and M. J. Matarić, "Learning from history for behavior-based mobile robots in non-stationary conditions," *Autonomous Robots and Machine Learning* AR:5, ML:31, 1998, AR:335-354, ML:141-167.
48. [Mitchell97] T. M. Mitchell, *Machine Learning*, McGraw-Hill, 1997.
49. [Moscato03] P. Moscato and C. Cotta, "A gentle introduction to memetic algorithms," In F. Glover and G. Kochenberger (Eds.), *Handbook of Metaheuristics*, pp. 105-144. Kluwer Academic Publishers, Boston MA, 2003.
50. [Moscato92] P. Moscato and M. Norman, "A 'Memetic' approach for the travelling salesman problem – implementation of a computational ecology for combinatorial optimisation on message-passing systems," In *Proceeding of the International Conference on Parallel Computing and Transputer Applications*, IOS Press, Amsterdam, 1992.
51. [Ng99] A. Y. Ng, D. Harada, and S. Russell, "Policy invariance under reward transformations: theory and application to reward shaping," *Proc. 16th International Conference on Machine Learning*, Morgan Kaufmann, San Francisco, CA, 1999.
52. [Nili01] M. Nili Ahmadabadi and E. Nakano, "A constrain and move approach to distributed object manipulation," *IEEE Transaction on Robotics and Automation*, vol. 17, no. 2, 2001, pp. 157-172.
53. [Nili02] M. Nili Ahmadabadi and M. Asadpour, "Expertness based cooperative Q-learning," *IEEE Transaction on Systems, Man, and Cybernetics – Part B: Cybernetics*, vol. 32, no. 1, 2002, pp. 66-76.
54. [Nolfi00] S. Nolfi and D. Floreano, *Evolutionary Robotics: The Biology, Intelligence, and Technology of Self-Organizing Machines*, MIT Press, Cambridge, MA, 2000.
55. [Ong04] Y.S. Ong and A. J. Keane, "Meta-lamarckian learning in memetic algorithms," *IEEE Transaction on Evolutionary Computation*, vol. 8, no. 2, pp. 99-110, 2004.
56. [Parker98] L. Parker, "ALLIANCE: an architecture for fault-tolerant multi-robot cooperation," *IEEE Transactions on Robotics and Automation*, 14 (2), 1998, pp. 220-240.
57. [Peshkin01] L. Peshkin, *Reinforcement Learning by Policy Search*, Ph.D. Dissertation, MIT, 2001.
58. [Precup00] D. Precup, *Temporal Abstraction in Reinforcement Learning*, Ph.D. Dissertation, *University of Massachusetts*, Amherst, MA, 2000.
59. [Puterman94] M. L. Puterman, *Markov Decision Process*, John Wiley and Sons, New York, 1994.

60. [Radcliffe94] N. J. Radcliffe and P. D. Surry, "Formal memetic algorithm," In T. Fogarty (ed.), *Evolutionary Computing: AISB workshop*, vol. 865 of Lecture Notes in Computer Science, pp. 1-16, Springer-Verlag, Berlin, 1994.
61. [Robbins51] H. Robbins and S. Monro, "A stochastic approximation method," *Annals of Mathematical Statistics* 22, 1951, pp. 400-407.
62. [Shoham04] Y. Shoham and R. Powers, "On the agenda(s) of research on multi-agent learning," *AAAI Fall Symposium on Artificial Multi-Agent Learning*, 2004.
63. [Sing94] S. Singh, R. Yee, "An upper bound on the loss from approximate optimal-value functions," *Machine Learning*, 16(3), pp. 227-233, 1994.
64. [Singh00] S. Singh, T. Jaakkola, M. Littman, and C. Szepesvari, "Convergence results for single-step on-policy reinforcement-learning algorithms," *Machine Learning*, vol. 38, no. 3, pp. 287-308, 2000.
65. [Singh04] S. Singh, A. G. Barto, S. Singh, N. Chentanez, "Intrinsically motivated reinforcement learning," In *Advances in Neural Information Processing (NIPS)*, 18, 2004.
66. [Sutton88] R. S. Sutton, "Learning to predict by the method of temporal differences," *Machine Learning*, 3(1), 1988, pp. 9-44.
67. [Sutton98] R. S. Sutton, and A. G. Barto, *Reinforcement Learning: An Introduction*, MIT Press, Cambridge, MA, 1998.
68. [Sutton99] R. S. Sutton, D. Precup, and S. Singh, "Between mdps and semi-mdps: a framework for temporal abstraction in reinforcement learning," *Artificial Intelligence*, 112, 1999, pp. 181-211.
69. [Thrun93] S. Thrun and A. Schwartz, "Issues in using function approximation for reinforcement learning," In *Proceedings of the 4th Connectionist Models Summer School*, 1993.
70. [Togelius04] J. Togelius, "Evolution of a subsumption architecture neurocontroller," *Journal of Intelligent and Fuzzy Systems*, 15:1, pp. 15-20, 2004.
71. [Urzelai98] J. Urzelai, D. Floreano, M. Dorigo, and M. Colombetti, "Incremental robot shaping," *Connection Science*, 10, 341-360, 1998.
72. [Vidyasagar03] M. Vidyasagar, *Learning and Generalization: with Applications to Neural Networks – 2nd edition*, Springer-Verlag, 2003.
73. [Wang96] Z. D. Wang, E. Nakano, and T. Matsukawa, "Realizing cooperative object manipulation using multiple behavior-based robots," in: *Proc. IEEE/RSJ Int. Conf. Intelligent Robots and Systems*, vol. 1, Osaka, Japan, 1996, pp. 310–317.
74. [Watkins89] C. J. C. H. Watkins, *Learning from Delayed Rewards*, Ph.D. Dissertation, Cambridge University, England, 1989.
75. [Watkins92] C. J. C. H. Watkins and P. Dayan, Q-learning, *Machine Learning*, 8(3), 1992, pp. 279-292.
76. [Williams94] R. J. Williams and L. C. Baird, "Tight performance bounds on greedy policies based on imperfect value functions," in: *Proc. 10th Yale Workshop on Adaptive and Learning Systems*, 1994.

77. [Zou04] P. Zou, Z. Zhou, G. Chen, and X. Yao, "A novel memetic algorithm with random multi-local-search: a case study of tsp," In Proceedings of the 2004 *Congress on Evolutionary Computation* (CEC04), pp. 2335-2340, 2004.

ضمیمه الف) واژه‌نامه فارسی-انگلیسی

از آن‌جا که بسیاری از واژه‌های تخصصی به کار رفته در این پایان‌نامه هنوز ترجمه‌ی مناسب و مقبول‌ای ندارند، لازم دیدم فهرستی از واژه‌هایی را که در طول متن به کار برده‌ام در ضمیمه‌ای جداگانه بیاورم. با وجود آن‌که هر کجا واژه‌ای جدید به کار رفته است معادل آن نیز در زیرنویس آمده، اما فهرست‌ای جداگانه بدون شک خواندن این پایان‌نامه را راحت‌تر می‌کند. پیشاپیش لازم است بگویم که نویسنده خود اعتقاد دارد برای بعضی از این واژه‌ها ترجمه‌ی مناسب‌تری نیز ممکن است.

Partial Observable	به طور جزئی رویت‌پذیر	Perception	ادراک
Return	پاداش بازگشته	Value	ارزش
Primitive	پایه‌ای	Incremental	اضافه‌شونده
Bottom up	پایین به بالا	Temporally Extended Action	اعمال گسترش-زمانی یافته
Coverage	پوشش‌دادن	Partition	افراز
Sample Complexity	پیچیدگی نمونه	Incremental	افزایشی
Decompose	تجزیه	Memetic Algorithm	الگوریتم ممیتیک
Excite	تحریک	Simulated Annealing	انجماد شبیه‌سازی‌شده
Progress Estimator	تخمین‌زن پیش‌رفت	Supervised	با سرپرست
Crossover	تزیوج	Recursive	بازگشتی
Projection	تصویر کردن	Representation	بازنمایی
Emergence	تظاهر	Planning	برنامه‌ریزی
Trial	تعداد آزمون‌ها	Dynamic Programming	برنامه‌ریزی پویا
Wall Following	تعقیب دیوار	Genetic Programming	برنامه‌ریزی ژنتیکی
Generalization	تعمیم	Bernoulli	برنولی
Commutativity	تعویض‌پذیری	Fork Lifting Mechanism	بلند کردن جسم با استفاده از مکانیزم چنگالی
General Function Approximator	تقریب‌زن عمومی		
Cultural Diversity	تنوع فرهنگی		

Behaviorism	رفتارگرایی	State Transition	توزیع احتمال انتقال
Observable	رویت‌پذیر	Probability	حالت
Tilt Angle	زاویه‌ی کج‌شدگی	Distributivity	توزیع‌پذیری
Landmark	زمین‌نشان	Distributed	توزیع‌شده
Subroutine	زیربرنامه	Baldwinian	تکامل بالدوینی
Subgoal	زیرهدف	Evolution	تکامل لامارکی
Subtask	زیر-وظیفه	Lamarckian	تکامل مصنوعی
genotype	ژنوتایپ	Evolution	تکامل هم‌کارانه‌ی
Fitness	سازگاری	Artificial	رفتارها
Suppress	سرکوب	Evolution of	تک‌گام
Bumper	سنسورهای برخوردی	Behaviors	جبر بول
Sonar	سونار	Single Step	جدول ارجاعی
Policy	سیاست	Boolean Algebra	جستجوی تصادفی
Policy-based	سیاست‌محور	Lookup Table	جستجوی کامل فضا
Cooperative	سیستم‌های چندعامله‌ی	Stochastic Search	جسم‌دار
Multi-Agent	هم‌کار	Exhaustive	جمع‌شونده بودن
Systems		Search	جهت‌ده اولیه
Reinforcement	سیگنال تقویت	Embodied	جهش
Signal	سیگنال تقویت	Contraction	چندگام
Shaped	شکل‌دهی‌شده	Bias	حافظه‌محور
Reinforcement	شرکت‌پذیری	Mutation	حریصانه
Signal	عامل	Multi Step	خبرگی
Associativity	عامل هوش‌مند	Memory-based	خنثی
Agents	عمل	Greedy	خودخواه
Intelligent Agent	عمل پایه‌ای	Expertness	خوشه‌بندی
Action	عمل پوچ	Neutral	در محیط قرار گرفته
Primitive Action	عینی	Selfish	دو به دو مجزا
No Action	غیرقابل پی‌گیری	Clustering	ذهنی
Objective	فاکتور فراموشی	Situated	راکتیو
Intractable	فرهنگ	Mutually	
Forgetting Factor		Exclusive	
Culture		Subjective	
		Reactive	

Scalable	مقیاس‌پذیر	Representation Space	فضای بازنمایی
Scaled	مقیاس‌شده	Hypothesis Space	فضای فرضیه
Meme	مم	Problem Space	فضای مساله
Monte Carlo	مونت کارلو	Activated	فعال شده
First Visit Monte Carlo Method	مونت کارلو اولین - ملاقات	Reusability	قابلیت استفاده مجدد
McCallum	مک کالوم	Function	کارکرد
Non-Stationary	ناایستا	Functionalism	کارکردگرایی
No Action	ناعمل	Complete	کامل
Invariance	ناوردا	Exploration	کاوش
Instance	نمونه	Global	کلی
Hopfield	هافیلد	Controlling	کنترل‌کننده
Premature Convergence	هم‌گرایی زودهنگام	Choice Node	گره‌های انتخاب
Reinforcement Learning	یادگیری تقویتی	Deceptive	گمراه‌کننده
Hierarchical Reinforcement Learning	یادگیری تقویتی سلسله‌مراتبی	Compliance	گیره
Contribution	یاری‌بهر	Layered	لایه‌ای
Monolithic	یک‌پارچه	Complement	متمم
		Moving Average	متوسط‌گیر لغزان
		Abstract	مجرد
		Actuator	محرك
		Local	محلی
		Environment	محیط
		Genetic Pool	مخزن ژنتیکی
		Meme Pool	مخزن مم
		Combinatorial Optimization Problem	مساله‌ی بهینه‌سازی ترکیبیاتی
		Purely Parallel SubSumption Architecture (PPSSA)	معماری Subsumption کاملاً موازی
		Behavior-based Architectures	معماری‌های رفتارگرا

ضمیمه ب) واژه‌نامه انگلیسی-فارسی

Abstract	مجرد	Controlling	کنترل کننده
Action	عمل	Cooperative	تکامل هم کارانه‌ی
Activated	فعال شده	Evolution of Behaviors	رفتارها
Actuator	محرك	Cooperative Multi-Agent Systems	سیستم‌های چندعامله‌ی هم کار
Agents	عامل	Coverage	پوشش دادن
Artificial Evolution	تکامل مصنوعی	Crossover	تزیج
Associativity	شرکت پذیری	Cultural Diversity	تنوع فرهنگی
Baldwinian Evolution	تکامل بالدوینی	Culture	فرهنگ
Behavior-based Architectures	معماری‌های رفتارگرا	Deceptive	گمراه کننده
Behaviorism	رفتارگرایی	Decompose	تجزیه
Bernoulli	برنولی	Distributed	توزیع شده
Bias	جهت ده اولیه	Distributivity	توزیع پذیری
Boltzman Action Selection	انتخاب عمل بولتزمان	Dynamic Programming	برنامه ریزی پویا
Boolean Algebra	جبر بول	Embodied	جسم دار
Bottom up	پایین به بالا	Emergence	تظاهر
Bumper	سنسورهای برخوردی	Environment	محیط
Choice Node	گره‌های انتخاب	Excite	تحریک
Clustering	خوشه بندی	Exhaustive Search	جستجوی کامل فضا
Combinatorial Optimization Problem	مساله‌ی بهینه سازی ترکیبیاتی	Expertness	خبرگی
Commutativity	تعویض پذیری	Exploration	کاوش
Complement	متمم	First Visit Monte Carlo Method	مونت کارلو اولین - ملاقات
Complete	کامل	Fitness	سازگاری
Compliance	گیره	Forgetting Factor	فاکتور فراموشی
Contraction	جمع شونده بودن	Fork Lifting Mechanism	بلند کردن جسم با استفاده از مکانیزم
Contribution	یاری بهر		چنگالی

Function	کارکرد	Memetic Algorithm	الگوریتم ممیتیک
Functionalism	کارکردگرایی	Memory-based	حافظه‌محور
General Function Approximator	تقریب‌زن عمومی	Monolithic	یک‌پارچه
Generalization	تعمیم	Monte Carlo	مونت کارلو
Genetic Pool	مخزن ژنتیکی	Moving Average	متوسط‌گیر لغزان
Genetic Programming	برنامه‌ریزی ژنتیکی	Multi Step	چندگام
Genotype	ژنوتایپ	Mutation	جهش
Global	کلی	Mutually Exclusive	دو به دو مجزا
Greedy	حریصانه	Neutral	خنثی
Hierarchical Reinforcement Learning	یادگیری تقویتی سلسله‌مراتبی	No Action	ناعمل، عمل پوچ
Hopfield	هاپفیلد	Non-Stationary	ناایستا
Hypothesis Space	فضای فرضیه	Objective	عینی
Incremental	اضافه‌شونده	Observable	رویت‌پذیر
Incremental	افزایشی	Partial Observable	به طور جزئی رویت‌پذیر
Instance	نمونه	Partition	افراز
Intelligent Agent	عامل هوش‌مند	Perception	ادراک
Intractable	غیرقابل پی‌گیری	Perturbation	اختلال
Invariance	ناوردا	Planning	برنامه‌ریزی
Lamarckian Evolution	تکامل لامارکی	Policy	سیاست
Landmark	زمین‌نشان	Policy-based	سیاست‌محور
Layered	لایه‌ای	Premature Convergence	هم‌گرایی زودهنگام
Local	محلی	Primitive	پایه‌ای
Lookup Table	جدول ارجاعی	Primitive Action	عمل پایه‌ای
McCallum	مک‌کالوم	Problem Space	فضای مساله
Meme	مم	Progress Estimator	تخمین‌زن پیش‌رفت
Meme Pool	مخزن مم	Projection	تصویر کردن
Memetic Algorithm	الگوریتم ممیتیک	Purely Parallel SubSumption Architecture (PPSSA)	معماری Subsumption کاملاً موازی

Reactive	راکتیو	Situated	در محیط قرار گرفته
Recursive	بازگشتی	Sonar	سونار
Reinforcement Learning	یادگیری تقویتی	State Transition Probability	توزیع احتمال انتقال حالت
Reinforcement Signal	سیگنال تقویت	Stochastic Search	جستجوی تصادفی
Representation	بازنمایی	Subgoal	زیرهدف
Representation Space	فضای بازنمایی	Subjective	ذهنی
Return	پاداش بازگشته	Subroutine	زیربرنامه
Reusability	قابلیت استفاده مجدد	Subtask	زیر-وظیفه
Sample	پیچیدگی نمونه	Supervised	با سرپرست
Complexity		Suppress	سرکوب
Scalable	مقیاس‌پذیر	Temporally Extended Action	اعمال گسترش-زمانی یافته
Scaled	مقیاس‌شده	Tilt Angle	زاویه‌ی کج‌شدگی
Selfish	خودخواه	Trial	تعداد آزمون‌ها
Shaped Reinforcement Signal	سیگنال تقویت شکل‌دهی‌شده	Value	ارزش
Simulated Annealing	انجماد شبیه‌سازی‌شده	Wall Following	تعقیب دیوار
Single Step	تک‌گام		

Abstract

The main goal of this research is the automatic design of behavior-based agents with hierarchical architecture. It is assumed that the designer can provide a reinforcement signal that reflects the behavior of the agent. Mathematically rigorous methods are provided to optimize this performance measure. I have divided the design problem to structure and behavior design sub-problems. Concurrent use of these methods can solve the whole design problem. In this research, I have provided learning-based and evolution-based methods for automating the design process.

Learning-based methods are offered for both behavior and structure design problems. To solve the structure learning problem, two different knowledge representations are introduced that by using them the agent's value is decomposed into simpler elements. Doing so, one can solve difficult issues such as credit assignment and value updating in this hierarchical architecture. A behavior learning method is provided that finds the suitable state-action mapping. By introducing a pseudo-action named NoAction, behaviors can cooperate with each other.

Co-evolution of behaviors is another method for behavior design. In contrast with the most other researches, behaviors are evolved separately and by using structure learning method, they are organized in the architecture. A new interpretation of cultural-inspired

method named Memetic Algorithm is offered that increases the performance of hybridization of evolution and learning.

The proposed methods are being tested in two different problems: multi-robot object lifting and an abstract problem. The results are competitive to human designed architectures.

In addition to the automatic design methods, a few other results are provided in this thesis. They are probabilistic analyzing of controlling probabilities of behaviors in our architecture and investigating the effect of reinforcement signal uncertainty on the agent's value function and its policy.



University of Tehran

Faculty of Engineering

Department of Electrical and Computer Engineering

Learning and Evolution in Hierarchical Behavior-based Systems

By: Amir massoud Farahmand

**Advisor:
Majid Nili Ahmadabadi**

**Co-Advisors:
Caro Lucas - Babak N. Araabi**

**A thesis submitted to the Graduate Studies Office
in partial fulfillment of the requirements for
the degree of M.Sc. in
Electrical Engineering**

August 2005