

پیوندگرایی

ترجمه: حسن صبوری مقدم
گروه روان‌شناسی - دانشگاه تبریز

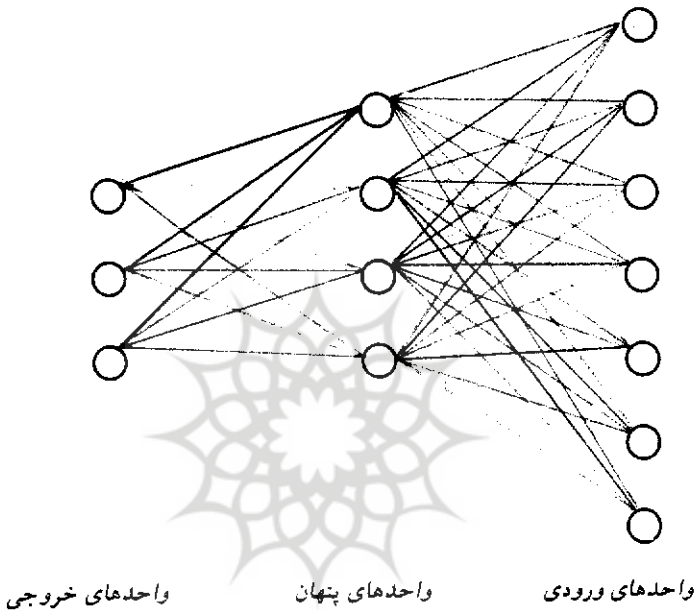
پیوندگرایی^۱ جنبشی است در علوم شناختی^۲ که امیدوار است توانایی‌های هوشی انسان را با استفاده از شبکه‌های عصبی مصنوعی^۳ (که تحت عنوان «شبکه‌های ارتباطی عصبی» یا «شبکه‌های عصبی» نیز شناخته می‌شود) تبیین نماید. شبکه‌های عصبی مدلهای ساده شده‌ای از مغز هستند که شمار زیادی از واحدها (در قیاس با سلولهای عصبی) را با یکدیگر ترکیب می‌کنند. ارتباط این واحدها دارای وزنهایی^۴ است که بیانگر قوت پیوند میان آنهاست. این وزنها مدلی است از تأثیرات سیناپسهای که یک سلول عصبی را به سلول عصبی دیگر پیوند می‌دهد. آزمایشها بر روی مدلهایی از این نوع ثابت کرده است که این مدلها توانایی یادگیری مهارتهایی همچون بازشناسی چهره، خواندن و کشف ساختار گرامری ساده را دارند.

علاقه فیلسوفان به پیوندگرایی بدین خاطر است که این دیدگاه حاوی جایگزینی برای نظریه کلاسیک در باب ذهن^۵ است؛ نظریه‌ای که وسیعاً معتقد است که ذهن چیزی از جنس یک کامپیوتر رقمی است که یک زبان نمادین را پردازش می‌کند. این که پارادایم پیوندگرا دقیقاً چگونه به چه میزان چالشی را برای کلاسیسم^۶ فراهم می‌کند یک موضوع بحث داغ در سالهای اخیر بوده است.

توصیفی از شبکه‌های عصبی

یک شبکه عصبی شامل شمار زیادی از واحدهای متصل به هم در یک الگویی از روابط است. واحدهای یک شبکه معمولاً در سه طبقه مجزا جای می‌گیرند: واحدهای ورودی^۷، که اطلاعاتی را که بایستی پردازش شوند دریافت می‌کنند؛ واحدهای خروجی^۸، که نتایج پردازش در آنها یافت می‌شود؛ و واحدهای میان ایندو که واحدهای پنهان^۹ نامیده می‌شود. اگر یک شبکه مدلی از کل سیستم عصبی انسان باشد، واحدهای

ورودی شبیه به نرونهاي حسي، واحدهاي خروجي شبیه به نرونهاي حرکتی، و واحدهاي پنهان شبیه به سایر نرونهاي سیستم عصبی می باشند. در اینجا شکل ساده از یک شبکه عصبی ساده آمده است:



هر واحد ورودی دارای یک ارزش فعالیت است که بخشی از محرکهای خارجی را برای شبکه باز نمایی می کند. یک واحد ورودی ارزش فعالیتش را به هر واحد پنهانی که با آن ارتباط دارد می فرستد. هر یک از واحدهای پنهان ارزش فعالیت خودشان را براساس ارزشهای فعالیتی که از واحدهای ورودی دریافت می کنند محاسبه می نمایند. سپس این سیگنال به سمت واحدهای خروجی یا سایر لایه های واحدهای پنهان فرستاده می شود. واحدهای پنهان نیز ارزش فعالیتشان را به همان طریق محاسبه می کنند و آنها را به واحدهای همجوارشان می فرستند. سرانجام اینکه سیگنال از واحدهای ورودی به تمام مسیرهای موجود در شبکه منتشر می شود تا ارزش فعالیت در واحدهای خروجی را مشخص سازد.

الگوی فعالیت موجود از یک شبکه توسط وزنها، یا قوت^۱ پیوندهای میان واحدها

تعیین می‌شود. وزن‌ها ممکن است مثبت یا منفی باشند. یک وزن منفی بیانگر اینست که فعال شدن واحد فرستنده باعث مهار واحد دریافت‌کننده می‌گردد. ارزش فعالیت برای هر واحد دریافت‌کننده بر طبق یک تابع فعالیت^{۱۱} ساده محاسبه می‌شود. جزئیات توابع فعالیت تغییر می‌کند، لیکن همه این تغییرات مطابق با همان نوع طرح اساسی است. تابع مشارکت همه واحدهای فرستنده را با هم جمع می‌زند، بدین ترتیب سهم هر واحد عبارت خواهد بود از وزن پیوند میان واحدهای فرستنده و گیرنده ضرب در ارزش فعالیت واحد فرستنده. معمولاً نتیجه بدست آمده بعداً تغییر می‌کند، برای مثال، با تعدیل مجموع فعالیت به ارزش میان صفر و یک و یا با رساندن فعالیت به صفر، در غیر اینصورت فعالیت به سطح آستانه خواهد رسید. پیوندگرایان بر این باورند که عملکرد شناختی می‌تواند توسط مجموعه واحدهایی که به این شکل عمل می‌کنند تبیین شود. در حالی که فرض بر اینست که همه واحدها تقریباً تابع فعالیت ساده یکسانی را محاسبه می‌کنند، توانایی‌های هوشمندانه انسان بایستی اساساً منوط به وضعیت وزنهای میان واحدها باشد.

نوع شبکه توصیف شده در بالا شبکه پیشخوراند^{۱۲} نامیده می‌شود، که در آن فعالیت مستقیماً از ورودی‌ها به واحدهای پنهان رفته و سپس بسوی واحدهای خروجی جریان می‌یابد. مدلهای واقعی تر مغز شامل چندین لایه واحد پنهان و روابط راجعه^{۱۳}، که سیگنالها را از سطوح بالاتر به سطوح پائین تر برمی گردانند، خواهند بود. چنین ارجاعی جهت تبیین خصوصیات شناختی چون حافظه کوتاه مدت ضروری است. در یک شبکه پیشخوراند، ارائه مکرر یک ورودی هر بار خروجی‌های یکسانی را تولید می‌کند، لیکن حتی ساده‌ترین ارگانیسم‌ها نیز به حضور تکراری محرکهای یکسان عادت می‌کنند (یا یاد می‌گیرند که این محرکها را نادیده بگیرند). پیوندگرایان مایلند که از پیوندهای راجعه اجتناب کنند، چرا که در مورد مسأله کلی آموزش^{۱۴} این شبکه‌ها کم می‌دانند. در عین حال، در خصوص شبکه‌های راجعه ساده، جایی که ارجاع کاملاً محدود است، آلان^{۱۵} و همکارانش (۱۹۹۱) پیشرفتهایی داشته‌اند.

یادگیری شبکه عصبی و انتشار روبه عقب

یکی از اهداف اصلی پژوهشهای پیوندگرا یافتن مجموعه درستی از وزن‌ها برای انطباق با

یک تکلیف فرضی است. خوشبختانه، الگوریتم‌های یادگیری^{۱۶} تهیه شده است که می‌تواند وزنهای درستی را برای انجام تعدادی از تکالیف محاسبه کند. (برای بازنگری موجود به کارهای هینتون^{۱۷} نگاه کنید، ۱۹۹۲). یکی از پرستفاده‌ترین روشهای آموزش، روش انتشار روبه عقب^{۱۸} نامیده می‌شود. برای استفاده از این روش نیاز به یک مجموعه آموزش شامل تعدادی نمونه ورودی و خروجی برای یک تکلیف فرضی است. برای مثال، اگر این تکلیف تمایز چهره مردان از زنان است، مجموعه آموزشی ممکن است حاوی تصاویری از چهره‌ها با مشخصه‌ای از جنسیت ترسیم شده در هر چهره باشد. شبکه‌ای که بتواند این تکلیف را یاد بگیرد ممکن است حاوی دو واحد خروجی (نماینده طبقات مردان و زنان) و تعدادی واحد ورودی، که یکی از آنها به درخشندگی هر پیکسل^{۱۹} (ناحیه بسیار کوچک) تصویر اختصاص یافته، باشد.

در شبکه‌ای که بایستی آموزش داده شود وزنها ابتدائاً به شکل ارزشهای تصادفی تعیین می‌شوند، و سپس شبکه به شکل مکرر در معرض عناصر مجموعه آموزشی قرار می‌گیرند. ارزشهای ورودی یک عضو در واحدهای ورودی قرار داده می‌شود، و خروجی شبکه با خروجی مطلوب برای این عضو مقایسه می‌شود. بنابراین، همه وزنها موجود در شبکه، در جهت اینکه ارزشهایی بیابند که به ارزشهای مورد نظر خروجی نزدیکتر باشد، تا حدی تعدیل می‌شوند. برای مثال، هنگامی که به واحدهای ورودی چهره مردی ارائه می‌شوند، وزنها به گونه‌ای تعدیل می‌شوند که ارزش واحد خروجی مرد افزایش یافته و ارزش واحد خروجی زن کاهش می‌یابد. بعد از این که این فرایند چندین مرتبه تکرار شد، ممکن است شبکه تولید خروجی مطلوب را برای هر ورودی موجود در مجموعه آموزشی یاد بگیرد. اگر آموزش بخوبی پیش برود، شبکه همچنین ممکن است تعمیم رفتار مطلوب را برای ورودی‌ها و خروجی‌هایی که در مجموعه آموزش قرار نداشته‌اند، یاد گرفته باشد. برای مثال، چنین شبکه‌ای ممکن است مهارت خوبی در تمایزگذاری میان تصاویر مردان از زنان، که قبلاً ارائه نشده، را نشان دهد.

آموزش شبکه‌ها به منظور مدلسازی جنبه‌هایی از هوش انسان هنر زیبایی است. موفقیت روش انتشار روبه عقب و سایر روشهای یادگیری پیوندگرا ممکن است منوط به سازگاری کاملاً دقیق الگوریتم و مجموعه آموزشی باشد. آموزش مشخصاً مستلزم

صدها از هزاران دور^{۲۰} سازگاری وزنی است. با توجه به محدودیت‌های کامپیوترهای فعلی که در دسترس محققین پیوندگرا است، آموزش یک شبکه برای اجرای یک تکلیف جالب ممکن است روزها یا حتی هفته‌ها وقت بگیرد. هنگامی که مدارهای موازی که اختصاصاً برای راه‌اندازی شبکه‌های عصبی طراحی شده‌اند وسیعاً در دسترس باشند، تعدادی از مشکلات حل می‌شود.

لیکن حتی در چنین مواردی نیز محدودیتهایی برای نظریه‌های یادگیری پیوندگرا باقی خواهد ماند که بایستی برای آنها فکری شود. انسانها (و تعداد کمی از حیوانات باهوش) توانایی یادگیری از وقایع منفرد را نشان می‌دهند؛ برای مثال حیوانی که غذایی را می‌خورد که باعث ناراحتی گوارشی‌اش می‌شود سعی خواهد کرد هرگز مجدداً آن را نخورد. تکنیکهای یادگیری پیوندگرا همچون انتشار روبه عقب ناتوان از تبیین این نوع یادگیری یک وهله‌ای^{۲۱} هستند.

نمونه‌هایی از آنچه که شبکه‌های عصبی می‌توانند انجام دهند

پیوندگرایان در اثبات قدرت شبکه‌های عصبی در تسلط یافتن بر تکالیفی شناختی پیشرفت چشمگیری نموده‌اند. در اینجا سه آزمایش کاملاً مشهور وجود دارد که پیوندگرایان را به این باور که شبکه‌های عصبی مدل‌های خوبی برای هوش انسان هستند، دلگرم کرده است. یکی از جالب توجه‌ترین این تلاشها کار سجنوفسکی^{۲۲} و روزنبرگ^{۲۳} (۱۹۸۷) بر روی شبکه‌ای است که می‌تواند متن انگلیسی را بخواند و شبکه گفتگو^{۲۴} نامیده شده است. مجموعه آموزشی برای شبکه گفتگو عبارت بود از بانک اطلاعاتی^{۲۵} بزرگی مشتمل بر یک متن انگلیسی همراه با خروجی آوایی مربوطه‌شان، که در یک کد مناسب برای استفاده در ارتباط با یک ترکیب کننده کلام بود. شنیدن نوارهای عملکرد شبکه گفتگو در مراحل متفاوت آموزش‌اش بسیار جالب هستند. در اولین خروجی سر و صدای اتفاقی است. در مرحله بعد، شبکه صداهایی را شبیه غان و غون کودکان تولید کرد، و بعداً این تولیدات تبدیل به گفتار دوتایی نه چندان واضح شد (گفتاری که شبیه به صدای کسانی بود که کلمات انگلیسی را تقلید می‌کنند). در پایان آموزش، شبکه گفتگو تکلیف تلفظ متنی را که به او داده شده بود بخوبی انجام داد. علاوه بر این، این توانایی به متنی که در مجموعه آموزشی ارائه نشده بود بخوبی تعمیم پیدا کرد.

مدل پیوندگرایی اولیه موثر دیگر شبکه‌ای بود که توسط رامله‌هارت^{۲۶} و مک کلند^{۲۷} (۱۹۸۶) تولید شد. این مدل برای پیش‌بینی زمان گذشته افعال انگلیسی آموزش داده شد. این تکلیف کار جالبی است، به این خاطر که گرچه اکثر افعال در انگلیسی (افعال با قاعده) زمان گذشته‌شان با اضافه کردن پسوند «ed» تولید می‌شود، تعدادی از افعال نیز بی‌قاعده هستند (مثل *go/ went* ، *come / came* ، *was/ is*). ابتدا به شبکه مجموعه زیادی از افعال بی‌قاعده آموزش داده شد، و سپس ۴۶۰ فعل که عمدتاً با قاعده بودند آموزش داده شد. شبکه زمان گذشته این ۴۶۰ فعل را در تقریباً ۲۰۰ دور آموزش یاد گرفت، و آن را به شکل کاملاً مناسبی به افعالی که در مجموعه آموزشی نبودند تعمیم داد. این شبکه حتی درک درستی از قواعدی که بایستی در میان افعال بی‌قاعده یافت شود، نشان داد (مثل *fly/ flew* ، *blow/ blew* ، *build/ built* ، *send/ sent*). در جریان آموزش، همان‌طور که سیستم در معرض یک مجموعه آموزشی که بیشتر حاوی افعال با قاعده قرار گرفت، گرایش بیشتری به قاعده‌مندی نشان داد، یعنی، اشکال فعلی با قاعده و بی‌قاعده را با هم قاطی کرد: (*break / broke* به جای *break/ broke*). این نقص با آموزش بیشتر تصحیح شد. جالب است خاطر نشان شود که همین گرایش به قاعده‌سازی افراطی در طی یادگیری زبان در کودکان دیده می‌شود. هر چند در این مورد که آیا مدل رامله‌هارت و مک کلند مدل خوبی از چگونگی یادگیری و پردازش جزء آخر فعل در زبان‌آموزی در انسانها هست یا خیر، بحث‌های داغی وجود دارد. برای مثال، پینکر^{۲۸} و پرینس^{۲۹} (۱۹۸۰) خاطر نشان کرده‌اند که این مدل در تعمیم به بعضی افعال با قاعده جدید ضعیف عمل می‌کند. آنها معتقدند که این ضعف نشانه‌ای از شکستی اساسی در مدل‌های پیوندگراست. شبکه‌ها ممکن است در ساختن الگوهای تداعی و مشابه‌سازی خوب باشند، لیکن آنها دارای محدودیتهای بنیادی در مهارت یافتن در قاعده‌های کلی همچون ساختن زمان گذشته افعال با قاعده می‌باشند. این ایرادها مسأله مهمی را برای مدل‌سازان پیوندگرا ایجاد می‌کند، اینکه آیا شبکه‌ها می‌توانند تکالیف شناختی عمده موجود در قواعد را تعمیم دهند. علی‌رغم ایرادهای پینکر و پرینس، تعدادی از پیوندگرایان معتقدند که روش تعمیم‌دهی درست هنوز امکان دارد (نیکلاسون^{۳۰} و ون‌گلدر^{۳۱}، ۱۹۹۴).

کارالمن (۱۹۹۱) در مورد شبکه‌هایی که می‌توانند ساختار گرامری را درک کنند

برای این مسأله که آیا شبکه‌های عصبی می‌توانند تسلط بر قواعد را یاد بگیرند یا نه، دلالت‌های مهمی مهمی داشت. ال‌من یک شبکه راجعه ساده را برای پیش‌بینی کلمه بعدی در یک مجموعه بزرگی از جملات انگلیسی آموزش داد. جملات متشکل بودند از یک دایره لغات ساده ۲۳ کلمه‌ای و با استفاده از مجموعه کوچکی از قواعد دستور زبان انگلیسی ساخته شده بودند. این دستور زبان، اگرچه ساده، آزمون سختی را برای آگاهی زبانی مطرح کرد. این دستور همزمان با آن‌که مستلزم مطابقت فعل و اسم اصلی جمله بود امکان ساخت بندهای نسبی نامحدودی را فراهم کرد. مثل جمله‌ای که در زیر می‌آید.

هر انسانی که بدنال سگهایی است که گربه‌هایی را دنبال می‌کنند... می‌دود.

فعل «می‌دود» بایستی با ضمیر خود «انسان» همخوان باشد، علی‌رغم مداخله اسمهای جمع (سگها)، «گربه‌ها» که ممکن است باعث شوند فعل «می‌دوند» انتخاب شود. یکی از ویژگیهای مهم مدلی ال‌من استفاده از ارتباطات راجعه است. در این ارتباطات، ارزشهای واحدهای پنهان در مجموعه‌ای از آنچه که واحدهای زمینه‌ای خوانده می‌شوند، ذخیره می‌شود تا برای دور بعدی پردازش به عقب به سطح ورودی فرستاده شوند. این حلقه روبه عقب، یعنی از واحدهای پنهان به لایه‌های ورودی، شکل اولیه‌ای از حافظه توالی کلمات موجود در جمله ورودی را برای شبکه فراهم می‌کند. شبکه‌های ال‌من ساختار گرامری جمله‌ای را که در مجموعه آموزشی نیست درک می‌کنند. عملکرد شبکه در مورد نحو^{۳۲} به طریق زیر مورد سنجش قرار گرفت. این بخش عبارت بود از پیش‌بینی کلمه بعدی در یک جمله انگلیسی، که البته تکلیف غیرممکنی است. در عین حال، این شبکه‌ها حداقل بوسیله این سنجش، در این کار موفق شدند. در این روش، در یک نقطه فرضی از جمله ورودی، بایستی واحدهای خروجی برای کلماتی که از لحاظ گرامری می‌توانند ادامه آن جمله باشند فعال شوند و برای سایر کلمات نافع‌ال گردند. بعد از آموزش متمرکز شدید، ال‌من توانست شبکه‌هایی تولید کند که در مورد جملاتی که حتی در مجموعه آموزشی نیامده بودند، عملکرد عالی نشان دهند. اگرچه این عملکرد مؤثر بود، لیکن هنوز برای دستیابی به شبکه‌های آموزش دیده‌ای که بتوانند پردازش زبانی داشته باشند راه طولانی در پیش است. علاوه بر این، در مورد اهمیت نتایج ال‌من تردیدهایی ایجاد شده است. برای مثال، مارکوس^{۳۳} (۲۰۰۱، ۱۹۹۸)

عنوان می‌کند که شبکه‌های *المن* قادر به تعمیم این عملکرد به جملاتی با کلمات تازه نیستند. او مدعی است که این امر نشان می‌دهد که مدل‌های پیوندگرا صرفاً نمونه‌ها را تداعی می‌کنند، و قادر به تسلط درست بر قواعد انتزاعی نیستند.

نقاط قوت و ضعف مدل شبکه‌های عصبی

فیلسوفان به این خاطر به شبکه‌های عصبی علاقه‌مندند که ممکن است آنها چارچوب جدیدی برای فهم ماهیت ذهن و ارتباطش با مغز فراهم کنند (راملهارت و مک کلند، ۱۹۸۶، فصل ۱). به نظر می‌رسد مدل‌های پیوندگرا بخوبی به آنچه که ما در مورد عصب‌شناسی می‌دانیم نزدیک شده‌اند. مغز در واقع یک شبکه نرونی است، که از تعدادی واحدهای متراکم (سلول‌های عصبی) و روابطشان (سیناپسها) شکل گرفته است. علاوه بر این، ویژگی‌های متعددی از مدل‌های شبکه‌های بعدی نشان می‌دهد که پیوندگرایی ممکن است بطور خاص تصویر کاملی از ماهیت پردازش شناختی ارائه دهد. شبکه‌های عصبی انعطاف‌پذیری زیادی در مقابل چالش‌های موجود در دنیای واقعی نشان می‌دهند. ورودی شلوغ یا خرابی واحدها سبب تنزل خفیفی در عملکرد می‌شود. معیذا پاسخ شبکه متناسب است، هرچند که صحت عملکرد تا حدی کاسته می‌شود. برعکس، شلوغی و فقدان شدت جریان برق در کامپیوترهای کلاسیک بطور مشخص منجر به نقصان شدیدی می‌گردد. شبکه‌های عصبی همچنین مشخصاً برای مشکلاتی که نیاز به حل تعدادی از محدودیتهای ناجور موجود در خطوط موازی دارند بخوبی سازگار شده‌اند. شواهد پژوهشی کافی در زمینه هوش مصنوعی^{۳۴} وجود دارد مبنی بر اینکه تکالیف شناختی چون بازشناسی شئی، طرح‌ریزی، و حتی حرکت هماهنگ، مشکلاتی از این نوع را نشان می‌دهند. اگر چه سیستم‌های کلاسیک قادر به رفع محدودیت‌های متعددی هستند، پیوندگرایان عنوان می‌کنند که مدل‌های شبکه عصبی حاوی مکانیسم‌های عصبی بیشتری برای برخورد با چنین مشکلاتی می‌باشند.

در طول قرن‌ها، فیلسوفان تلاش کرده‌اند بفهمند مفاهیم ما چگونه تعریف می‌شوند. در حال حاضر عموماً اعتراف می‌شود که سعی در مشخص کردن تصورات عادی در قالب شروط لازم و کافی محکوم به شکست است. تقریباً برای هر تعریف پیشنهادی استثنائاتی وجود دارد. برای مثال، شخصی ممکن است یک ببر را یک گربه سیاه و

نارنجی بزرگ تعریف کند. لیکن درباره ببر سفید چه فرض خواهد کرد؟ فیلسوفان و روانشناسان شناختی عنوان می‌کنند که طبقه‌بندی مفاهیم به شکل انعطاف‌پذیرتری تعیین می‌شود، برای مثال، یک طبقه‌بندی ممکن است از طریق توجه به شباهت یا قرابت به یک نمونه نخستین^{۳۵} صورت بگیرد. ظاهراً مدل‌های پیوندگرا بطور خاصی برای تطبیق تصورات درجه‌بندی شده از عضویت مقوله‌ای از این نوع خیلی مناسب می‌باشند. شبکه‌ها می‌توانند تصدیق الگوهای آماری ثابت را یاد بگیرند. الگوهایی که با خشک و سخت شدن قواعد ارائه‌شان خیلی سخت خواهد شد. پیوندگرایی متعهد تبیین انعطاف‌پذیری و بینش موجود در هوش انسان است، و این تعهد را از طریق استفاده از روش‌هایی که نمی‌توانند بسادگی در قالب رد اصول آزاد بیان شوند، به انجام می‌رساند (هورگان^{۳۶} و تینسون^{۳۷}، ۱۹۸۹، ۱۹۹۰)، بنابراین از شکنندگی ناشی از آشکال استاندارد بازنمایی‌های نمادین^{۳۸} اجتناب می‌کنند.

علیرغم این چهره فریبنده، ضعف‌هایی در مدل‌های پیوندگرا وجود دارد که در ادامه به آنها اشاره می‌شود. اولاً، اکثر پژوهش‌های مربوط به شبکه‌های عصبی از بعضی از خصوصیات جالب و احتمالی مغز دور شده‌اند. برای مثال، پیوندگرایان معمولاً نه تلاشی برای تصریح مدل‌های متنوع انواع سلول‌های عصبی مغزی می‌کنند، و نه تلاشی در مورد ناقل‌های عصبی و هورمون‌ها انجام می‌دهند. علاوه بر این، بعید به نظر می‌رسد که مغز حاوی نوعی روابط معکوس باشد، روابطی که در یادگیری انتشار روبه عقب نیاز است، و به نظر می‌رسد که شمار بیش از حد تکرارهایی که برای چنین آموزش‌هایی نیاز است دور از واقعیت باشد. اگر قرار است مدل‌های پیوندگرای متقاعدکننده‌ای در مورد پردازش شناختی انسان ساخته شود توجه به این موضوعات احتمالاً ضروری خواهد بود. یک ایراد جدی‌تر دیگری نیز مطرح است. ایرادی که بخصوص از سوی طرفداران کلاسیسم مطرح می‌شود اینست که شبکه‌های عصبی بطور خاص برای نوعی پردازش پایه که تصور می‌شود زیان، استدلال و اشکال عالی‌تر تفکر را تقویت می‌کند، مناسب نیستند. ما در بحث نظام‌مدار بودن این موضوع را بررسی خواهیم کرد.

بازنمایی پیوندگرا^{۳۹}

مدل‌های پیوندگرا پارادایم جدیدی را برای فهم اینکه چگونه اطلاعات ممکن است در

مغز بازنمایی شود، بدست می‌دهد. یک ایده ساده، لیکن گمراه کننده، اینست که یک سلول عصبی ساده (با دسته کوچکی از سلولهای عصبی) ممکن است به بازنمایی هر چیزی که مغز نیاز دارد ضبط کند، اختصاص یابد. برای مثال، ما ممکن است تصور کنیم که یک سلول عصبی اختصاص به مادر بزرگ دارد و هنگامی که ما درباره مادر بزرگمان فکر می‌کنیم این سلول شلیک می‌کند. لیکن، چنین بازنمایی کانونی^{۴۱} غیر متحمل است و شواهد خوبی وجود دارد مبنی بر این که تصور کردن مادر بزرگمان مستلزم الگوهای پیچیده‌ای از فعالیت توزیع شده در قسمتهای نسبتاً وسیعی از مغز می‌باشد.

جالب است خاطر نشان شود که بازنمایی‌های توزیع شده^{۴۱} در واحدهای پنهان، بیشتر از بازنمایی‌های کانونی، محصولات طبیعی روشهای آموزش پیوندگرا هستند. برای نمونه مادامی که شبکه گفتگو متن را پردازش می‌کند الگوهای فعالیت در واحدهای پنهان ظاهر می‌شوند. تجزیه و تحلیل نشان می‌دهد که این شبکه برای بازنمایی حروف بی صدا و صدا دار با ایجاد یک واحد فعال برای حروف صدا دار و واحدی دیگر برای حروف بی صدا عمل نمی‌کند بلکه بیشتر این عمل را با رشد دو الگوی فعالیت مشخصاً متفاوت در کل واحدهای پنهان انجام می‌دهد.

با توجه به انتظارات شکل گرفته از تجربه مان با بازنمایی کانونی، به نظر می‌رسد فهم بازنمایی توزیع شده هم تازه باشد و هم مشکل. لیکن این تکنیک مزایای عمده‌ای را در بردارد. برای مثال، هنگامی که قسمتهایی از مدل خراب یا اضافه بار^{۴۲} می‌شوند، بازنمایی‌های توزیع شده (بی شباهت به نمادهای ذخیره شده در محلهای حافظه ثابت جداگانه) نسبتاً بخوبی حفظ می‌شوند. مهم‌تر این که، در حالی که بازنمایی‌ها بیش از آن که در شلیک واحدهای منفرد کدگذاری می‌شوند در الگوها^{۴۳} کدگذاری شوند، و روابط میان بازنمایی‌ها در شباهت‌ها و تفاوت‌های میان این الگوها کدگذاری می‌شود. چنان که خصوصیات درونی این بازنمایی حاوی اطلاعاتی در مورد آن چیزی است که آنها بازنمایی می‌کنند (کلارک^{۴۴} ۱۹۹۳، صفحه ۱۹). برعکس، بازنمایی کانونی قراردادی است. هیچ ویژگی ذاتی بازنمایی (شلیک یک واحد) روابطش را با سایر نمادها تعیین نمی‌کند. این چهره گویا از بازنمایی‌های توزیع شده متعهد حل یک معمای فلسفی در مورد معناست. در یک طرحواره^{۴۵} بازنمایانه نمادین، همه بازنمایی‌ها فارغ از اتمها (همانند کلمات در یک زبان) ساخته می‌شوند. معنای رشته‌های پیچیده نمادین ممکن

است خارج از عناصر تشکیل دهنده شان تعریف شوند، لیکن چه چیزی معنای این اتمها را ثابت نگه می‌دارد؟

طرحهای بازنمایانه پیوندگرا با صرف نظر کردن ساده از اتمها یک راه حل نهایی برای این معما بدست می‌دهند. هر بازنمایی توزیع شده الگویی از فعالیت در طول همه واحدهاست، چنان که هیچ راه ثابت شده‌ای برای تمایز میان بازنمایی‌های ساده از بازنمایی‌های پیچیده وجود ندارد. یقیناً، بازنمایی‌ها فارغ از فعالیت واحدهای منفرد بنا می‌شوند. لیکن هیچ یک از این «اتمها» برای یک نماد کدگذاری نمی‌شوند. بازنمایی‌ها زیر نمادین^{۴۶} هستند به این معنا که تجزیه در اجزای ترکیب کننده شان و رای سطح نمادین رها می‌شود.

ماهیت زیر نمادین بازنمایی توزیع شده راه تازه‌ای برای تصور پردازش اطلاعات در مغز بدست می‌دهد. اگر ما فعالیت هر سلول عصبی را با یک شماره مدلسازی کنیم، بنابراین فعالیت کل مغز می‌تواند بوسیله یک بردار خیلی بزرگ (یالیستی) از شماره‌ها، هر شماره را برای یک سلول عصبی، بدست آید. هم ورودی مغز از سیستم‌های حسی و هم خروجی‌اش به سلولهای عصبی - عضلانی می‌توانند به عنوان بردارهایی از این نوع در نظر گرفته شود. بطوری که مغز به یک پردازشگر بردار تبدیل می‌شود و مشکل

۲۴۰

روان‌شناسی تبدیل خواهد شد به پرسشهایی درباره این که عملکرد کدام بردارها کدام جنبه شناختی انسان را تبیین می‌کنند.

بازنمایی زیر نمادین دلالت‌های جالبی برای فرضیه کلاسیکی دارد مبنی بر این که مغز بایستی حاوی بازنمایی‌های نمادینی باشد که مشابه جملات یک زبان است. این ایده، که غالباً تر زبان تفکر^{۴۷} (یا LOT) نامیده می‌شود ممکن است توسط بازنمایی پیوندگرا به چالش کشیده شود. نمی‌توان بسادگی کاملاً آنچیزی را که تر LOT تبیین می‌کند گفت، لیکن وان گلدر^{۴۸} (۱۹۹۰) نشانه موثر و قابل قبولی را برای تعیین زمانی که بایستی گفته شود مغز حاوی بازنمای‌های شبه جمله^{۴۹} است، ارائه می‌دهد. این هنگامی است که بازنمایی یک مورد نشانه‌گذاری شده بوسیله نشانه‌های آن بازنمایی ساخته می‌شود. برای مثال، اگر من بنویسم «جان مری را دوست دارد» من به موجب ساختمان جمله نوشته شده عناصر زیر را خواهم داشت: «جان»، «دوست دارد» و «مری».

بازنمایی‌های توزیع شده‌ای برای اظهارات پیچیده‌ای چون «جان مری را دوست دارد» می‌تواند ساخته شود که حاوی هیچ بازنمایی صریحی از قسمت‌هایش نیست (اسمولسکی^{۵۰}، ۱۹۹۱). اطلاعات مربوط به اجزای تشکیل دهنده می‌تواند از بازنمایی‌ها استخراج شود، لیکن مدل‌های شبکه عصبی نیازی به استخراج صریح خود این اطلاعات جهت پردازش صحیح آن ندارند (چالمرز^{۵۱}، ۱۹۹۰). این مسأله نشان می‌دهد مدل‌های شبکه عصبی به عنوان نقطه مقابل این ایده که زبان تفکر شرط لازم برای فرایندهای شناختی انسان است، عمل می‌کنند. هرچند، این موضوع هنوز یک بحث داغ در این حوزه است (فودور^{۵۲}، ۱۹۹۷).

شکل مباحثه میان پیوندگرایان و کلاسیک‌ها

طی سی سال گذشته دیدگاه کلاسیک غالب بوده است، مبنی بر این که فرایندهای شناختی انسان (حداقل شناخت‌های عالی‌تر) را با محاسبه نمادین در کامپیوترهای رقمی قابل مقایسه می‌داند. در این تبیین کلاسیک، اطلاعات بوسیله رشته‌هایی از نمادها بازنمایی می‌شود، درست همانطوری که ما داده‌ها را در حافظه کامپیوتر یا در برگه‌های کاغذ بازنمایی می‌کنیم. از سوی دیگر، پیوندگرایان ادعا می‌کنند که اطلاعات به شکل غیر نمادین در وزن‌ها، یا قوت پیوندها، میان واحدهای یک شبکه عصبی ذخیره می‌شود. طرفداران کلاسیک اعتقاد دارند که شناخت به پردازش رقمی شباهت دارد، جایی که رشته‌ها برطبق دستورات یک برنامه (نمادین) به شکل متوالی تولید می‌شود. پیوندگرایان پردازش ذهنی را به عنوان تکامل تدریجی گام به گام و پویای فعالیت در یک شبکه عصبی در نظر می‌گیرند، و فعالیت هر واحد منوط به قوت پیوند و فعالیت واحد همجوارش، مطابق با تابع فعالیت، خواهد بود.

علی‌الظاهر، این دیدگاه‌ها خیلی متفاوت به نظر می‌رسند. هر چند تعدادی از پیوندگرایان کارشان را به عنوان چالشی با کلاسیسم در نظر نمی‌گیرند و تعدادی نیز آشکارا از دیدگاه کلاسیک حمایت می‌کنند. بطوری که این گروه، که در پی ایجاد یک همگرایی میان دو پارادایم هستند، پیوندگرایان اجرایی^{۵۳} خوانده شده‌اند. آنها معتقدند که شبکه مغزی بر یک پردازشگر نمادین دلالت دارد. بدرستی، ذهن یک شبکه عصبی است، لیکن این شبکه همچنین در سطح انتزاعی‌تر و بالاتر از آنچه که توصیف شده، یک پردازشگر نمادین می‌باشد. بدین ترتیب بر طبق نظر پیوندگرایان اجرایی، نقش

پژوهشهای پیوندگرا کشف چگونگی دستگاه مورد نیاز پردازش نمادین است که می‌تواند از مواد شبکه عصبی ساخته شود، و نقش پردازش کلاسیک می‌تواند به شرح شبکه عصبی کاهش یابد.

به هر حال، تعدادی از پیوندگرایان بر روی نکته اجرایی این دیدگاه پافشاری می‌کنند. چنین پیوندگرایان رادیکالی مدعی‌اند که تصور پردازش در مورد کارکرد ذهن تصوری است نامناسب. انتقاد آنها این است که نظریه کلاسیک تبیین ضعیفی از کاهش خفیف عملکرد، بازنمایی کل‌نگرانه داده‌ها، تعمیم همزمان، درک زمینه، و سایر ویژگیهای هوشی انسان که در مدلهایشان گنجانده شده بدست می‌دهند. شکسته برنامه‌نویسی کلاسیک برای جور شدن با انعطاف‌پذیری و کارآمدی شناختهای انسان نشانه‌ای از نیاز به پارادایمی جدید در علوم شناختی است. بطوری که پیوندگرایان رادیکال پردازش نمادین را برای همیشه از علوم شناختی کنار گذاشته‌اند.

بحث نظام‌دار بودن^{۵۴}

نکته عمده مجادله در ادبیات فلسفی پیوندگرایی بایستی در ارتباط با این موضوع باشد که آیا پیوندگرایان پارادایمی جدید و قابل دوام برای فهم ذهن بدست می‌دهند یا خیر؟ یک انتقاد اینست که مدلهای پیوندگرا فقط برای پردازش تداعی‌ها مناسب‌اند. لیکن تکالیفی همچون زبان و استدلال نمی‌توانند صرفاً بوسیله روشهای تداعی‌گر به انجام رسند، همچنانکه احتمالاً پیوندگرایان نیز با عملکرد مدلهای کلاسیکی در تبیین این توانایی‌های شناختی عالی موافق نیستند. در عین حال، مادامی که شبکه‌هایی بتوانند ساخته شوند که مدارهای یک کامپیوتر را تقلید نمایند، اثبات این موضوع که هرکاری که پردازشگرهای نمادین می‌توانند انجام دهند شبکه‌های عصبی نیز می‌توانند انجام دهند امر ساده‌ایست. بنابراین این ایراد نمی‌تواند وارد باشد که مدلهای پیوندگرا نمی‌توانند تبیینی برای کارکردهای عالی شناختی باشند، بلکه بایستی گفت آنها فقط در صورتی می‌توانند این کار انجام دهند که ابزارهای پردازش نمادین تکمیل شده باشد. بدین ترتیب پیوندگرایی اجرایی ممکن است موفق باشد، لیکن پیوندگرایان رادیکال هرگز قادر به تبیین ذهن نخواهند شد.

مقاله اصلی مورد استشهدا فودور و پلیشین^{۵۵} (۱۹۸۸) بحثی از این دست را مطرح

می‌سازد. آنها ویژگی از هوش انسان را شناسایی کردند که نظامدار بودن خوانده شده است، و آنها احساس می‌کنند که پیوندگرایان قادر به تبیین این ویژگی نیستند. نظامدار بودن زبان به این حقیقت باز می‌گردد که توانایی تولید / فهم / تفکر بعضی جملات ذاتاً به توانایی تولید / فهم / تفکر سایر ساختارهای مرتبط باز می‌گردد. برای مثال، هیچ فردی با تسلطی بر زبان انگلیسی که زبان «جان مری را دوست دارد» را می‌فهمد نمی‌تواند در فهم «مری جان را دوست دارد» مشکل داشته باشد. از دیدگاه کلاسیک، رابطه میان این دو توانایی می‌تواند بسادگی تبیین شود. با این فرض که افراد مسلط به زبان انگلیسی اجزای تشکیل دهنده جمله (جان)، «دوست دارد» و «مری» «جان مری را دوست دارد» را بازنمایی می‌کنند و معنایش را از معانی این تشکیل دهنده‌ها استنباط می‌کنند. اگر این درست باشد، بنابراین جمله تازه‌ای مثل «مری جان دوست دارد» نیز می‌تواند به عنوان نمونه دیگری از همین فرایند نمادین تبیین شود. به شیوه‌ای مشابه، پردازش نمادین تبیینی خواهد بود برای نظامدار بودن استدلال، یادگیری، و تفکر. بدین ترتیب نظام‌دار بودن تبیین می‌کند که چرا هیچ فردی وجود ندارد که قادر به نتیجه‌گیری P از $(p \ \& \ q)$ باشد، لیکن از نتیجه‌گیری P از $P \ \& \ Q$ ناتوان باشد، چرا هیچ فردی وجود ندارد که قادر به یادگیری تقدم یک مکعب قرمز بر یک مربع سبز باشد و ناتوان از یادگیری تقدم یک مکعب سبز بر یک مربع قرمز باشد، و چرا هیچ فردی وجود ندارد که بتواند تصور کند که مری جان را دوست دارد ولی نتواند تصور کند که جان مری را دوست دارد.

فودور و مک لوقلین^{۵۶} (۱۹۹۰) با جزئیات بیشتری عنوان می‌کنند که پیوندگرایان تبیینی برای نظامدار بودن ندارند. اگرچه مدل‌های پیوندگرا می‌تواند آموزش ببینند نظامدار باشند، آنها همچنین می‌توانند، برای مثال، برای بازشناسی «جان مری را دوست دارد» بدون توانایی بازشناسی «مری جان را دوست دارد» آموزش ببینند. زمانی که پیوندگرایی تضمینی برای نظامدار بودن نمی‌کند، برای نظامدار بودن گسترده شناخت انسان نیز تبیینی ندارد. نظامدار بودن ممکن است در معماری‌های پیوندگرا وجود داشته باشد، لیکن این نظامدار بودن موجود بیشتر از یک اتفاق تصادفی نیست. راه‌حل کلاسیک تاحدی بهتر است، چرا که در مدل‌های کلاسیکی نظامدار بودن عمدتاً ناشی از آزادی است.

این اتهام که شبکه‌های پیوندگرا در تبیین نظامدار بودن ناتوان هستند توجه زیادی را

جلب کرده است. نکته‌ای که عمدتاً برای رد این اتهام مورد استشهاد قرار گرفته است (آیزاوا^{۵۷}، ۱۹۹۷؛ ماتیوس^{۵۸}، ۱۹۹۷، هادلی^{۵۹} ۱۹۹۷ ب) اینست که معماری‌های کلاسیکی نیز در تبیین نظام‌دار بودن بهتر عمل نمی‌کنند. همچنین مدلهای کلاسیکی وجود دارند که می‌توانند برای بازشناسی جمله «جان مری را دوست دارد» بدون توانایی بازشناسی جمله «مری جان را دوست دارد» برنامه‌ریزی شده باشند. نکته اینست که نه استفاده از معماری پیوندگرا و نه استفاده از معماری کلاسیکی به تنهایی قادر به تبیین نظام‌دار بودن نیستند. در هر دو معماری، فرضهای بعدی در مورد ماهیت پردازش بایستی به گونه‌ای باشد که بتواند جمله «جان مری را دوست دارد» را نیز پردازش کند. بخشی در این مورد بایستی اشاره‌ای به درخواست فودور و مک‌قلیان داشته باشد مبنی بر این که نظام‌دار بودن بایستی به عنوان یک موضوع عادی، یعنی به عنوان یک موضوع قانون طبیعی تبیین شود. انتقاد علیه پیوندگرایان اینست که هر چند که ممکن است آنها سیستم‌هایی را بکار گیرند که نظام‌دار بودن را نشان بدهد، لیکن این سیستمها نظام‌دار بودن را تبیین نخواهند کرد مگر در مدلهایشان این موضوع به عنوان یک ضرورت عادی دنبال شود. در عین حال، تقاضا برای ضرورت عادی تقاضای بزرگی است، تقاضایی که هیچ یک از معماریهای کلاسیک به روشنی نمی‌توانند به آن دست یابند، پس تنها تدبیر برای رفع این ایراد در این راستا خواهد بود که درخواست تبیین نظام‌دار بودن برای موردی که معماری‌های کلاسیک می‌توانند آنها احراز کنند و برای معمارهای پیوندگرایی که نمی‌توانند آنها احراز کنند کاهش بیابد. متعزداً مورد قانع‌کننده‌ای از این نوع بایستی ساخته شود.

پیوندگرایی و شباهت معنایی

یکی از جاذبه‌های بازنمایی‌های توزیع شده در مدلهای پیوندگرا اینست که آنها برای مشکل تعیین معنای حالت‌های مغز راه‌حلی پیشنهاد می‌کنند. ایده آنها این است که اطلاعات معنایی در شباهتها و تفاوت‌های میان الگوهای فعالیت همراه ابعاد متفاوت فعالیت عصبی ضبط شود. از این طریق، ویژگیهای مشابه فعالیت‌های نرونی ویژگیهای ذاتی را که معنا ایجاد می‌کند بدست می‌دهد، هر چند، فودور و لپور^{۶۰} (۱۹۹۲، فصل ۶)

تبیین‌های پایه‌ای شباهت^{۶۱} را در دو زمینه به چالش کشیده‌اند. اولین مشکل اینست که از قرار معلوم تعداد سلولهای عصبی و روابط میان این سلولها در مغز انسانها بطور معنی‌داری با هم متفاوتند. گرچه در تعریف سنجش شباهتها در دو شبکه‌ای که حاوی تعداد یکسانی از واحدهاست امر ساده‌ایست، ولی هنگامی که معماری‌های پایه دو شبکه متفاوت است دیدن این که چگونه این امر می‌تواند انجام شود کار مشکل‌تری است. دومین مشکل که فودور و لپور ذکر می‌کنند اینست که گرچه سنجش شباهتها برای معانی می‌تواند بطور موفقیت‌آمیزی با مهارت اجرا شود، لیکن آنها برای احراز یک مورد ضروری که لازمه یک نظریه معنی است ناکافی خواهند بود.

چارچلند^{۶۲} (۱۹۹۸) نشان داد که اولین ایراد ایندو می‌تواند پاسخ داده شود. او با استشهد به کار لاکسو^{۶۳} و کوتزل^{۶۴} (۲۰۰۰) توضیح داد که چگونه سنجش شباهتها میان الگوی فعالیت‌ها در شبکه‌هایی با ساختارهای کاملاً متفاوت می‌تواند تعریف شود. غیر از این مسأله نیز لاکسو و کوتزل با سنجشهایی صورت گرفته نشان دادند شبکه‌هایی با ساختارهای متفاوت که در تکلیف یکسانی آموزش دیده بودند الگوهای فعالیتی کاملاً مشابهی را نشان می‌دهند. این امر این امید را ایجاد می‌کند که بخشهایی که به لحاظ تجربی بخوبی تعریف شده باشند در خصوص شباهت مفاهیم و افکار در افراد مختلف می‌تواند ساخته شود.

از سوی دیگر، طرح یک نظریه سنتی در مورد معنی مبتنی بر شباهت موانع بزرگی را پیش‌رو خواهد داشت (فودور و لپور، ۱۹۹۹)، برای چنین نظریه‌ای تعیین شرایط درست جملات با توجه به تحلیلی از معنای قسمت‌هایشان نیاز خواهد بود، و روشن نیست که شباهت به تنهایی بتواند تفکیک ثابتی، بگونه‌ای که یک نظریه استاندارد می‌طلبد، را انجام دهد. هر چند، اکثر پیوندگرایانی که مروج شباهت به عنوان تبیین اساسی معنا هستند تعدادی از این پیش‌فرضهای نظریه‌های استاندارد را رد می‌کنند. آنها در حالیکه هنوز به داده‌هایی در مورد توانایی‌های زبانی انسان وفادارند، امیدوارند که با یک کار دقیق مقدماتی دیگر این پیش‌فرضها را یار د کنند و یا تغییر داده و اصلاح نمایند. با فرض فقدان یک نظریه خوب در مورد معنا در پارادایم‌های سنتی یا پیوندگرا، بهتر است که این مسأله به پژوهشهای آتی واگذار شود.

پیوندگرایی و حذف روان‌شناسی عامیانه

کاربرد مهم دیگر پژوهشهای پیوندگرا در بحث فلسفی مربوط به ذهن توجه به وضعیت روان‌شناسی عامیانه^{۶۵} [قومی] است. روان‌شناسی عامیانه ساختار مفهومی است که ما بطور خودانگیخته برای فهم و پیش‌بینی رفتار انسان بکار می‌بریم. برای مثال، دانستن اینکه جان میل به یک نوشیدنی دارد و اینکه او می‌داند در یخچال نوشیدنی وجود دارد به ما اجازه می‌دهد که چرایی رفتن او را به آشپزخانه تبیین کنیم. چنین دانشی قطعاً منوط به این خواهد بود که تصور کنیم دیگران امیال و اهدافی دارند و طرحهایی برای ارضاء آنها و باورهایی برای هدایت این طرحها خواهند داشت. این ایده که افراد دارای باورها، طرحها و امیالی می‌باشند یک مسأله پیش پا افتاده در زندگی عادیست؛ لیکن آیا این ایده توصیف درستی از آنچه که واقعاً در مغز اتفاق می‌افتد بدست می‌دهد؟

مدافعین روان‌شناسی عامیانه عنوان می‌کنند که این روان‌شناسی آنقدر خوب و مناسب است که نمی‌تواند اشتباه باشد (فودور، ۱۹۸۸، فصل ۱). ما برای تأیید صحت یک نظریه چه چیزی بیشتر از این می‌خواهیم که این نظریه چارچوب ضروری برای گفتگوی موفقیت‌آمیز با دیگران فراهم کند؟ از سوی دیگر، حذف‌گرایان^{۶۶} پاسخ خواهند داد که استفاده وسیع و مفید از یک طرحواره مفهومی دلیلی بر صحت‌اش نیست (چارچلند، ۱۹۸۹، فصل ۱). منجمان عهد باستان تصور می‌کردند که افلاک آسمانی در اداره نظامشان مفید (حتی اساسی) است، لیکن اکنون ما می‌دانیم که هیچ فلک آسمانی وجود ندارد. از نقطه نظر حذف‌گرایان، وفاداری به روان‌شناسی عامیانه، مثل وفاداری به طب عامیانه (ارسطویی) در راه پیشرفت علمی مانع ایجاد می‌کند. یک روان‌شناسی قابل رشد ممکن است نیاز به یک حرکت تند انقلابی در بنیادهای مفهومی‌اش داشته باشد، حرکتی که در ماشینهای کوانتومی یافت می‌شود.

حذف‌گرایان به این خاطر به پیوندگرایی علاقه‌مند هستند که این دیدگاه مدعی ایجاد بنیاد مفهومی است که ممکن است جانشین روان‌شناسی عامیانه باشد. برای مثال، رامسی^{۶۷} و همکارانش (۱۹۹۱) عنوان کرده‌اند که شبکه‌های پیشخوراند معین تکالیف شناختی ساده‌ای را نشان می‌دهند که می‌توانند بدون بکارگیری خصوصیات که می‌توانستند مطابق باورها، امیال و طرحها باشند عمل کنند. با فرض اینکه چنین شبکه‌هایی به چگونگی کار مغز وفادار هستند، مفاهیم روان‌شناسی عامیانه خوراکی

بهتر از آنچه که افلاک آسمانی برای این شبکه‌ها فراهم می‌کنند، ایجاد نخواهد نمود. اینکه آیا مدل‌های پیوندگرا روان‌شناسی عامیانه را در این راه تحلیل برده و تخریب خواهند کرد یا خیر هنوز جای بحث است. در مورد این ادعا که مدل‌های پیوندگرا از نتایج حذف‌گرایی حمایت می‌کنند، دو ایراد اصلی وجود دارد. یک ایراد اینست که مدل‌های استفاده شده توسط رامسی و همکارانش شبکه‌های پیشخوراندی هستند که در تبیین بعضی از اساسی‌ترین ویژگی‌ها شناختی همچون حافظه کوتاه مدت بیش از حد ضعف دارند. رامسی و همکارانش نشان نداده‌اند که باورها و امیال در طبقه‌ای از شبکه‌های مناسب برای شناخت‌های انسان حذف شوند. خط دومی از چالش‌های انتقادی مدعی است که حتی در شبکه‌های پیشخوراندی مورد بحث نیز حذف ویژگی‌های متطبق با باورها و امیال ضروری است (فون‌اشارت^{۶۸}، زیر چاپ).

با عدم توافق بر سر ماهیت روان‌شناسی عامیانه، مسأله پیچیده‌تر می‌شود. تعدادی از فیلسوفان باورها و امیال فرض شده روان‌شناسی عامیانه را به عنوان حالت‌های مغزی با محتوی نمادین در نظر می‌گیرند. برای مثال، این تصور که در یخچال نوشیدنی وجود دارد، به عنوان یک وضعیت مغزی که حاوی نمادهایی در رابطه با نوشیدنی و یخچال است در نظر گرفته می‌شود. از نقطه نظر این دیدگاه، سرنوشت روان‌شناسی عامیانه قویاً به فرضیه پردازش نمادین گره خورده است. چنان که پیوندگرایان بتوانند ثابت کنند که پردازش مغزی اساساً غیرنمادین است، استنتاج‌های حذف‌گرایان دنبال خواهد شد. از سوی دیگر، تعدادی از فیلسوفان تصور نمی‌کنند که روان‌شناسی عامیانه اساساً نمادین باشد، و تعدادی حتی این ایده را که روان‌شناسی عامیانه بایستی به عنوان نظریه‌ای دارای اولویت در نظر گرفته شود، به چالش خواهند کشید. با این تصور، ایجاد رابطه‌ای میان نتایج پژوهش‌های پیوندگرا و زد روان‌شناسی عامیانه بسیار مشکل‌تر خواهد شد.

پی نوشتها

- 1- connectionsim
- 2- congitive science
- 3- artificial neural networks
- 4- weights
- 5- mind
- 6- classicism
- 7- input units
- 8- output units
- 9- hidden units
- 10- strength
- 11- feed forward
- 12- recurrent
- 13- activation function
- 14- traininy
- 15- Elman
- 16- learning algorithms
- 17- Hinton
- 18- backpropagation
- 19- pixel
- 20- training set
- 21- one shot
- 22- Sejnowski
- 23- Rosenberg
- 24- NET talk



- 25- data base
- 26- Rumelhart
- 27- Mcclelland
- 28- Pinker
- 29- Prince
- 30- Niklasson
- 31- Van Gelder
- 32- syntax
- 33- Marcus
- 34- artificial intelligence
- 35- prototype
- 36- Horgan
- 37- Tienson
- 38- symbolic representation
- 39- connectionist representation
- 40- local representation
- 41- distributed representation
- 42- overloaded
- 43- patterns
- 44- clark
- 45- scheme
- 46- sub - symbolic
- 47- language of thought
- 48- Van Gelder
- 49- sentence - like
- 50- Smolensky
- 51- Chalmers



- 52- Fodor
- 53- implementational connectionists
- 54- systematicity
- 54- pylyshyn
- 56- Mc Laughlin
- 57- Aizawa
- 58- Matthews
- 59- Matthews - Hadley
- 60- Lepore
- 61- similarity
- 62- Churchland
- 63- Laakso
- 64- Cottrell
- 65- folk psychology
- 66- eliminativists
- 67- Ramsey
- 68- Von Eckhart



شهرتگاه علوم انسانی و مطالعات فرهنگی
پرتال جامع علوم انسانی



پروہش گاہ علوم انسانی و مطالعات فرہنگی
پرتال جامع علوم انسانی