

فصلنامه پژوهش در مسائل تعلیم و تربیت

شماره پیاپی ۵۱ - دوره دوم ، شماره ۳۴ - بهار ۹۳

مقاله شماره ۱ - صفحات ۵ تا ۲۴

مدل سازی تشخیصی شناختی (CDM) با استفاده از نرم افزار R

دکترافشین افزلی^۱

چکیده

مدل های تشخیصی شناختی (CDM)، برای تعیین تسلط و یا عدم سلطه از افراد در صفات مختلف بر اساس یک ماتریس از ویژگی های از پیش تعیین شده توسعه یافته اند. بر خلاف نظریه سوال- پاسخ (IRT)، که بر تحلیل سطح سوال و نمرات کسب شده توسط آزمودنی ها متمرکز هستند، مدل های تشخیصی شناختی بر بازخوردهای غیررسمی و تشخیصی در مورد مهارتهای خاصی که باید بهبود یابند، برای دانش آموزان و معلمان تمرکز دارند. مدل های مختلفی برای سنجش تشخیصی شناختی طراحی شده اند که مدل دینا یکی از پرکاربرترین آنهاست. این مدل دو پارامتر حدس و لغزش را برای هر سوال به همراه ضرایب تشخیص و شاخص برازندگی سوال را برآورد می نماید. همچنین احتمال و فراوانی مورد انتظار هر یک از الگوهای صفات و همچنین احتمال کناری هر یک از صفات مکنون در این مدل محاسبه شده و الگوی صفات هر آزمودنی مشخص می گردد. نرم افزارهای مختلفی برای مدل های CDM پیشنهاد شده اند؛ در این مقاله مدل سازی CDM با استفاده از نرم افزار R به دلیل مزایایی که این نرم افزار دارد، به همراه یک نمونه پژوهشی با توضیح و تفسیر نتایج شامل پارامترهای سوال، احتمال کناری دستیابی به اهداف و همبستگی تتراکوریک بین اهداف ارائه شده است.

کلمات کلیدی: مدل های تشخیصی شناختی، مدل DINA، نرم افزار R

^۱ عضو هیئت علمی دانشگاه بوعلی سینا همدان <afzali.afshin@yahoo.com>

مقدمه

سنجش تشخیص شناختی نوعی از سنجش تربیتی است که به منظور اندازه‌گیری ساختارهای دانش و فرآیندهای مهارت‌های خاص در دانش‌آموزان با هدف فراهم کردن اطلاعات در رابطه با ضعف و قوت‌های آن‌ها، طراحی شده است (لاتیون و گریل، ۲۰۰۷) و هدف آن فراهم کردن اطلاعات معتبر آموزشی است تا با استفاده از آن‌ها معلم بتواند به‌طور مؤثر آموزش‌های بعدی و مکمل را برای پاسخ‌گویی به نیازهای یادگیری هر دانش‌آموز خاص طراحی کند.

در طول ۲۰ سال گذشته، علاقه متخصصین روانسنجی در زمینه مدل‌های تشخیصی شناختی^۱ (CDMs)، گسترش یافته است. این مدل‌ها شامل مدل‌هایی آماری هستند که به پژوهشگران اجازه می‌دهند که به‌صورت تجربی فرضیاتی را درباره طبیعت فرآیندهای پاسخدهی که آزمودنی‌ها در زمان پاسخگویی به یک سنجش یا سؤالات یک پرسشنامه استفاده می‌کنند، آزمون نمایند. اگر طرح جمع‌آوری داده‌ها و تئوری زیرساخت، کفایت مناسبی داشته باشد، اطلاعات تجربی دقیقی درباره مؤلفه‌های ذهنی مداخله‌کننده در فرآیندهای پاسخدهی و چگونگی تعامل میان این مؤلفه‌ها توسط این مدل‌ها قابل دستیابی است. همچنین، مدل‌های CDM یک طبقه‌بندی چند متغیره از آزمونی‌ها ایجاد می‌نماید که آن‌ها را براساس مؤلفه‌های ذهنی که زمینه‌ساز فرآیندهای پاسخدهی آن‌هاست، نیمرخ‌بندی می‌نماید، که این نیمرخ‌ها می‌تواند به منظور شناسایی مسیرهای اصلاحی برای شایستگی در همه مؤلفه‌ها، کمک کننده باشد. (راپ، ۲۰۰۷)

مدل‌های تشخیصی شناختی (CDM) مدل‌های متغیر مکنون چندبعدی تأییدی احتمالاتی با یک ساختار بارگذاری پیچیده هستند. این مدل‌ها برای مدلسازی متغیرهایی با پاسخ طبقه‌ای مناسب هستند و شامل متغیرهای پیش‌بینی‌کننده مکنون طبقه‌ای هستند که طبقات نهفته را تولید می‌نمایند. این مدل‌ها قادر به تولید تفسیرها و بازخوردهای چند ملاکی برای اهداف تشخیصی هستند که به یک تئوری زمینه‌ای - شناختی برای فرآیندهای پاسخ در اندازه ریز اشاره دارد (راپ، ۲۰۰۷).

با مدل‌سازی بر اساس شیوه‌های سنتی تئوری سؤال-پاسخ^۲ توانایی‌های آزمودنی‌ها در طول یک پیوستار مرتب می‌شوند و معمولاً یک نمره مقیاسی و یا یک رتبه درصدی به عنوان نمره گزارش شده، ارائه می‌شود،

¹ Cognitive Diagnostic Models

² Rupp

³ Item Response Theory

اما نتایج حاصل از نمره گذاری بر اساس مدل های تشخیصی شناختی (CDMs) متفاوت است، در این مدل ها آزمودنی ها پروفایل های مهارت چند بعدی را دریافت می کنند که بر اساس چیرگی یا عدم چیرگی در همه مهارت های موجود در آزمون طبقه بندی شده اند (دیبلو، راسوز و استات، ۲۰۰۷).

استفاده از روش های CDM در سنجش مزایایی دارد که در سایر روش ها وجود ندارد، در درجه نخست این که اکثر مدل های IRT، تک بعدی بودن آماری مجموعه داده ها را فرض نموده و آن را به عنوان شرط قبلی برای مدرج نمودن سؤالات و برآورد پارامترها نیاز دارند، به عبارت دیگر، در اکثر مدل ها، تک بعدی بودن به عنوان لازمه اصلی برای مشخص نمودن جایگاه آزمودنی ها در طی یک پیوستار فرضی، ضروری است. یکی از ویژگی های مهم CDM عدم نیاز به تک بعدی بودن است. به نظر می رسد تک بعدی بودن تا حدودی در بخش های تحصیلی مشکل ساز باشد چرا که تحقیقات نشان داده است که ابزارهای سنجش تحصیلی نوعاً مجموعه ای از صفات یا خرده مهارت ها را مورد توجه قرار می دهند که هر کدام از آن ها می تواند یک بعد آماری جداگانه را ایجاد نماید (آریا دوست، ۲۰۱۱).

هم چنین به عنوان یک از مدل های IRT چند بعدی، CDM یکی از روش های نسبتاً تازه طراحی شده در میان ارزشیاب های صفت مکنون است که هم زمان هم از روش های ریاضی برای برآورد پارامترها و هم از اصول روان شناسی شناختی برای متمایز نمودن افراد مسلط و غیر مسلط استفاده می نماید. (گیرل، سیوی و ژو، ۲۰۰۹) مدل های تشخیصی شناختی از این جهت ساخته شده اند تا از طریق اطلاعات هدفمندتری که به شکل پروفایل های نمره ارائه می دهند، نقص و محدودیت مدل های IRT را برطرف نمایند (دی لا توره، ۲۰۰۹).

شکلهای متنوعی از مدل های تشخیصی شناختی (DCMs) در تاریخچه سنجش مطرح شده است. بطور کلی این مدل ها انواع موقعیتهایی (مثل انواع سازه، پاسخ، و ابعاد) را پوشش می دهند که مورد علاقه پژوهشگران در روانسنجی و علوم شناختی و یادگیری می باشد. در حال حاضر پژوهشگران به دلیل گستردگی این مدل ها در تلاشند تا روابطی را بین این مدل ها بیابند و شباهتهایی را جهت یکسان سازی آنها بدست بیاورند. بعنوان مثال مدل DINA ی تعمیم یافته (دی لا توره، ۲۰۱۳)، مدل تشخیصی عمومی (وان داویر، ۲۰۰۵)، مدل تشخیصی شناختی طولی (هنسون، تمپلین و وایلز، ۲۰۰۹) اثبات نمودند که روند یکسان سازی مدل ها در

¹ DiBello, Roussos & Stout

² De la Torre

³ Von Davier

⁴ Henson., Templin & Willse

تاریخچه سنجش وجود دارد. این تلاشها منجر به محبوبیت فزاینده این مدلها شد و نیز موجب تحولی در بین پژوهشگران شد تا مدلهای CDM را برای محققان کاربردی ساده تر و امکانپذیرتر نمایند (یانگ سان^۱ و همکاران، ۲۰۱۱).

دو روش برای استفاده از سنجش تشخیصی شناختی برای سنجش مهارتها وجود دارد:

الف) تحلیل دادههای حاصل از یک سنجش موجود با استفاده از مدلهای پیچیدهتر مهارت محور به امید استخراج اطلاعات بیش تر نسبت به مقیاس سازیهای تک بعدی یا سایر روشهای موجود.

ب) طراحی یک آزمون از ابتدا برای یک هدف تشخیص مهارت محور. در طراحی یک سنجش تشخیصی مؤثر برآورده کردن اهداف سنجشی، انگیزه اصلی است. در اغلب کارهایی که در زمینهی مدلهای تشخیصی انجام شده و در ادبیات پژوهش یافت می شود از سنجش تشخیصی به عنوان یک روش پس از وقوع (راه اول) استفاده شده است (دیبلو و دیگران، ۲۰۰۷).

مدل DINA یک از پرکاربردترین مدلهای CDM در ادبیات این روشهاست. این مدل بدون توجه به تعداد صفتهای بررسی شده در هر ارزیابی به تنهایی برای آیتم نیازمند به تخمین دو پارامتر می باشد. بعلاوه زمانی که صفتهای موردنیاز برای یک آیتم دارای اهمیت برابر و مساوی هستند، مدل DINA مناسب بنظر می رسد. مطالعات مختلفی مدل DINA را بررسی نموده اند. یک بحث کلی از مدل DINA که شامل کاربردها و مدلهای طبقه بندی مرتبط و پنهانی آن است در کارهای دلاتوره و داگلاس^۲ (۲۰۰۸)، داگنون و فالماگن (۱۹۹۹)، هرتل (۱۹۸۹)، جانکر و سیجسما^۳ (۲۰۰۱)، مک ریدی و دایتون (۱۹۷۷)، و تاتسوکا (۲۰۰۲) قابل دستیابی می باشد.

مدل DINA شامل دو بخش می باشد: (۱) فرایند قطعی یا دوراز خطا و (۲) فرایند تصادفی. فرایند قطعی با بردار پاسخ پنهان (η_i) نشان داده می شود که با استفاده از بردار برآورد شده α_i محاسبه می شود؛ عبارت دیگر

$$\eta_i = \{\eta_{ij}\}$$

فرمول ۱

¹ Young-sun

² de la Torre & Douglas

³ Junker & sijtsma

$$\eta_{ij} = \prod_{k=1}^K \alpha_{ik}^{q_{jk}}$$

فرمول ۲

ارزش عددی ۱ به پاسخ پنهان نشاندهنده این است که آزمودنی دارای همه صفتهای مورد نیاز جهت حل آیت j می باشد در حالیکه ارزش عددی صفر به معنای آن است که آزمودنی یکی از صفتهای مورد نیاز برای آیت را دارا نیست، نه اینکه فاقد همه صفتهای k باشد. مدل DINA قسمت تصادفی موجود در این مدل را نیز پوشش می دهد زیرا پاسخهای دانش آموز اغلب دارای لغزش (خطاهای) غیرمنظم می باشد. قسمت دروازه ای «and» این مدل اسمش را از طبیعت ارتباطی η_{ij} بدست می آورد (دی لا تور ۲۰۰۹). همچنین دو پارامتر جنبه ی سروصدا را مشخص می کند: پارامترهای حدس و پارامترهای لغزش. مدل DINA پارامتر لغزش را با عنوان $s_j = P(X_{ij} = 0 | \eta_{ij} = 1)$ و پارامتر حدس را با عنوان $g_j = P(X_{ij} = 1 | \eta_{ij} = 0)$ برای آیت j مشخص می کند. برخلاف موقعیت غیرتصادفی دانش آموزانی که دارای همه صفتها نیستند و گاهی اوقات حدس^۱ می زنند (پارامتر حدس) و پاسخ صحیح به یک آیت می دهند و دانش آموزانی که دارای همه صفتها هستند و گاهی لغزش^۲ می کنند (پارامتر لغزش) و غلط به یک آیت پاسخ می دهند. مدل DINA با ترکیب این دو پارامتر این احتمال را محاسبه می کند که آزمودنی i چگونه آیت j را حل می نماید با توجه به این مسئله که مسیر مهارتهای i به صورت زیر فرمولبندی می شود

$$P_j(\underline{\alpha}_i) = P(X_{ij} = 1 | \underline{\alpha}_i) = g_j^{1-\eta_{ij}} (1 - s_j)^{\eta_{ij}}$$

فرمول ۳

ازاینرو برای ارائه پاسخ صحیح به یک آیت معادله شماره (۱) نیازمند آن است که یک آزمودنی همه صفتهای لازم برای حل آیت j را داشته باشد آنگونه که ماتریکس Q با اجتناب از خطا آنرا تعریف می کند و برای یک آزمودنی که فاقد صفت خاصی می باشد حدس صحیحی بزند. همچنین اگر هیچکدام از پارامترهای حدس و لغزش وجود ندارند، پاسخ آزمودنی کاملاً قطعی می شود و نتایج بر تعامل بین مسیر $\underline{\alpha}$ و بردار Q برای آیت متکی می باشد. (دی لاتور، ۲۰۰۹). تخمین تأثیرات هم پارامترهای حدس و هم پارامترهای لغزش می

¹ guessing

² slipping

تواند منجر به شناسایی یکسری روابط غیردقیق بین تسلط دانش آموز برصفت مورد نیاز و پاسخ دهی به یک آیتم گردد. درنتیجه مربیان و معلمان می توانند بفهمند که آیا یک صفت یا مهارت ویژه جهت دستیابی به یک سطح رتبه بندی معین بر یک دانش آموز یا گروهی از دانش آموزان تسلط دارد. جزئیات الگوریتم EM برای مدل DINA توسط دی لاتوره (۲۰۰۹) ارائه شده اند.

روش:

نرم افزارهای مختلفی به منظور انجام محاسبات و مدلسازی تشخیصی شناختی معرفی شده اند از جمله :

- ۱- نرم افزار R
- ۲- نرم افزار Mplus
- ۳- نرم افزار Apreggio
- ۴- نرم افزار OX
- ۵- نرم افزار FlexmIRT
- ۶- نرم افزار Matematica

با توجه به ویژگیهای مثبتی که نرم افزار R دارد از جمله اینکه یک نرم افزار رایگان و متن باز است و به راحتی قابل استفاده است و همچنین مدل‌های بنیادی CDM را پوشش می دهد، این مقاله به معرفی و آموزش مدلسازی تشخیصی شناختی با استفاده از این نرم افزار اختصاص یافته است. در ادامه پس از معرفی اجمالی این نرم افزار مراحل و گامهای اصلی به منظور انجام محاسبات مدلسازی تشخیصی شناختی با استفاده از این نرم افزار ارائه می گردد. سپس در بخش یافته ها نتایج یک مثال پژوهشی با استفاده از داده های شبیه سازی شده به همراه تحلیل خرو جی ها ارائه می گردد.

نرم افزار R به همراه برخی عملکردهای پایه‌ای داندود می شود که آماده استفاده هستند، با این حال به منظور بهره گیری از بسیاری از عملکردهای نیاز به نصب بسته‌هایی^۱ است که کاربر را قادر به استفاده‌های اختصاصی از این نرم افزار می کنند.

به منظور استفاده از R برای - مدلسازی تشخیص شناختی نیاز به نصب بسته‌ی نرم افزاری CDM می باشد.

البته بسته‌ای نیز به نام NPCD نیز برای مدل‌های ناپارامتری تشخیص شناختی وجود دارد

بسته CDM فرآیند برآورد پارامترها برای دو مدل اصلی CDM یعنی مدل‌های DINA و DINO (دلاتوره و داگلاس، ۲۰۰۴، جانکرو سیجسما، ۲۰۰۱، تمپلین و هنسون ۲۰۰۶)، مدل تعمیم یافته‌ی DINA برای

¹ Packages

صفات دو ارزشی (GDINA، دلاتوره ۲۰۱۱) و صفات چند ارزشی (GDINA، چن^۱ و دلاتوره، ۲۰۱۳) مدل تشخیصی تعمیم یافته (GDM، ون داویر، ۲۰۰۸) و مدل گسترش یافته‌ی آن برای مدل طبقات متغیر مکنون چندبعدی IRT (بارتولوچی^۲، ۲۰۰۷) و همچنین ابزارهای لازم برای تحلیل داده‌ها تحت مدل‌های فوق را، فراهم می‌سازد. (روبیترز^۳، ۲۰۱۳)

همانگونه که گفته شد با توجه به کاربرد فراوان مدل DINA در این مقاله به این مدل پرداخته شده است، اطلاعات لازم برای استفاده از سایر مدل‌ها در راهنمای بسته‌ی CDM که توسط روبیترز و همکاران (۲۰۱۳) تدوین شده، قابل دستیابی است.

برای بارگذاری بسته CDM می‌توان از طریق منو اقدام نمود به این منظور باید از منوی Packages، گزینه load package را انتخاب نمود سپس در پنجره‌ای که باز می‌شود که شامل لیستی از تمامی بسته‌ها نصب شده است، در این لیست گزینه CDM انتخاب می‌شود. برای انجام تحلیل‌های CDM دو دسته داده مورد نیاز است:

الف) فایل داده‌ها ب) ماتریس Q

در هر بار استفاده از نرم‌افزار R ضروری است، مسیر کاری^۳ مطابق آدرس که فایل‌های فوق در آن ذخیره شده‌اند تغییر یا به‌طور مثال اگر فایل‌های فوق در دایرِی E و پوشه‌ی work ذخیره شده باشند با دستور زیر مسیر تغییر خواهد یافت.

Setwd ("E:/work")

همچنین این عمل را می‌توان از طریق منوی File و انتخاب گزینه‌ی Change dir... و در ادامه انتخاب مسیر مورد نظر انجام داد.

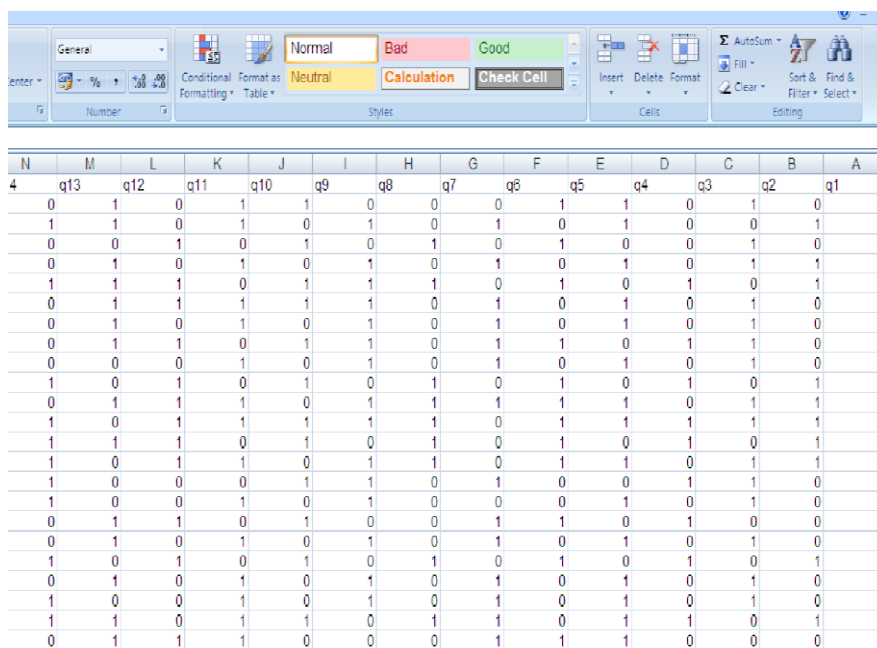
لازم به ذکر است که کلیه خروجی‌های نرم‌افزار R تا پایان استفاده از نرم‌افزار در این نوبت در آدرس فوق ذخیره خواهد شد. اما در هر بار استفاده از نرم‌افزار R ضروری است اقدامات فوق دوباره انجام شود. نرم‌افزار R به راحتی قابلیت استفاده از داده‌های ذخیره شده در نرم‌افزارهای معمول آماری از جمله SPSS و Excel را دارد. اما برای این کار باید داده‌ها را با فرمت .txt، در Excel و CSV، در SPSS ذخیره نمود. همانگونه که توضیح داده شد فایل داده‌ها شامل پاسخ آزمون‌ها در سطر و سؤالات در ستون‌های فایل است تصویر شماره بخشی از فایل داده‌های مورد استفاده در این پژوهش را نمایش می‌دهد، در این فایل سطر

¹ Chen

² Robitzsch

³ Working directory

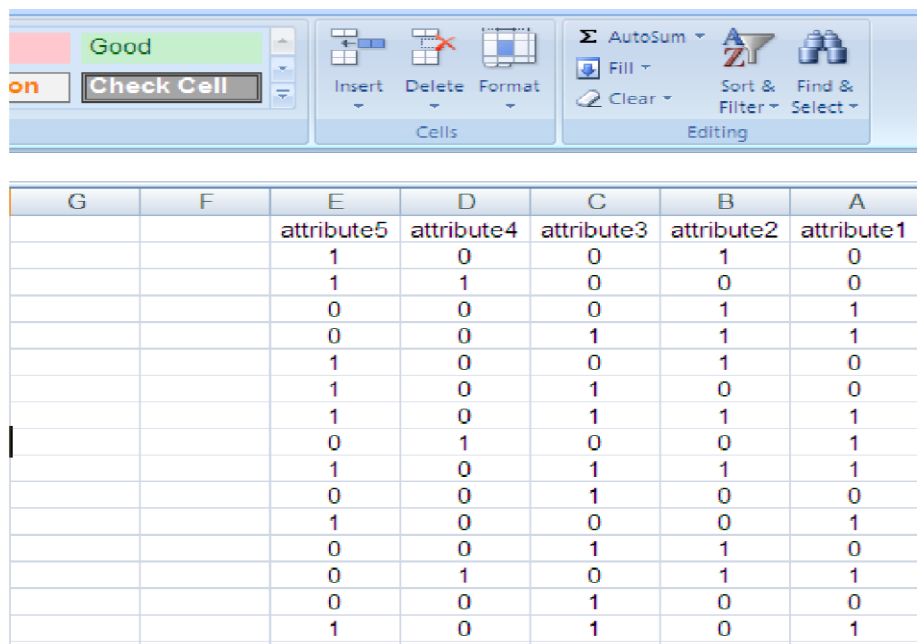
اول نامی است که به هر سوال اختصاص داده شده و سطرهاى بعدی پاسخ هر آزمودنی به هر سوال است (برای پاسخ درست و صفر برای پاسخ نادرست).



	N	M	L	K	J	I	H	G	F	E	D	C	B	A
q13	0	1	0	1	1	0	0	0	1	1	0	1	0	1
q12	1	1	0	1	0	1	0	1	0	1	0	0	1	1
q11	0	0	1	0	1	0	1	0	1	0	0	1	0	1
q10	0	1	0	1	0	1	0	1	0	1	0	1	1	0
q9	1	1	1	0	1	1	1	0	1	0	1	0	1	0
q8	0	1	1	1	1	1	0	1	0	1	0	1	0	1
q7	0	1	0	1	0	1	0	1	0	1	0	1	0	1
q6	0	1	1	0	1	1	0	1	1	1	0	1	0	1
q5	0	0	0	1	0	1	0	1	0	1	0	1	0	1
q4	1	0	1	0	1	0	1	0	1	0	1	0	1	0
q3	0	1	1	1	0	1	1	1	1	1	0	1	1	1
q2	1	0	1	1	1	1	1	0	1	1	1	1	1	0
q1	1	1	1	0	1	0	1	0	1	0	1	0	1	0
	1	0	1	0	1	0	1	0	1	0	1	0	1	0
	0	1	1	1	0	1	1	1	1	1	0	1	1	1
	1	0	1	1	0	1	1	0	1	1	0	1	1	0
	1	0	0	0	1	1	0	1	0	0	1	1	0	1
	1	0	0	1	0	1	0	0	0	1	0	1	0	0
	0	1	1	0	1	0	0	1	1	0	1	0	0	1
	0	1	0	1	0	1	0	1	0	1	0	1	0	1
	1	0	1	0	1	0	1	0	1	0	1	0	1	0
	0	1	0	1	0	1	0	1	0	1	0	1	0	1
	1	0	0	1	0	1	0	1	0	1	0	1	0	1
	1	1	0	1	1	0	1	1	0	1	1	0	1	1
	0	1	1	1	0	0	0	1	1	1	0	0	0	1

شکل ۱ فایل داده های ذخیره شده در Excel

در ماتریس Q بر اساس مدل شناختی که از قبل طراحی شده چنانچه سوالی به یک صفت مرتبط باشد به آن عدد ۱ در صورت عدم ارتباط عدد صفر اختصاص می یابد. در فایل ماتریس Q سؤالات در سطرها و صفتها در ستونها قرار می گیرند. شکل شماره ۲، فایل ماتریس Q را که در این مقاله استفاده شده است، نشان می دهد.



G	F	E	D	C	B	A
		attribute5	attribute4	attribute3	attribute2	attribute1
		1	0	0	1	0
		1	1	0	0	0
		0	0	0	1	1
		0	0	1	1	1
		1	0	0	1	0
		1	0	1	0	0
		1	0	1	1	1
		0	1	0	0	1
		1	0	1	1	1
		0	0	1	0	0
		1	0	0	0	1
		0	0	1	1	0
		0	1	0	1	1
		0	0	1	0	0
		1	0	1	0	1

شکل ۲ ماتریس Q شامل پنج صفت ذخیره شده در Excel

برای فراخوانی این فایل ها در نرم افزار R ، به طور مثال اگر فایل داده ها با نام Data1 و فایل ماتریس Q با نام Qmatrix و هر دو با پسوند txt. (از طریق نرم افزار Excel) در آدرسی که مسیر کاری R است ذخیره شده باشند برای خواندن این فایل ها به ترتیب از دستورات زیر استفاده می شود.

```
Datafile <- read.table ("Data1.txt , header = TRUE)
```

```
Qmtx<-read.table("Q matrix.txt , header = TRUE)
```

با دستورات فوق جدولی از داده ها تحت عنوان Data file و جدولی از ماتریس Q به نام Qmtx در نرم افزار R ذخیره می شود. عبارت TRUE در مقابل header به این معناست که سطر اول فایل داده ها، عنوان هر ستون است، در صورتی که فایل داده ها فاقد عنوان باشد واژه FALSE ذکر می شود. لازم به ذکر است که در هر بار استفاده از نرم افزار R برای کاربر روی داده های فوق باید جداول با دستورات فوق فراخوانی شود.

همچنین نرم افزار R می تواند به صورت غیرمستقیم و به کمک بسته نرم افزاری foreign داده های نرم افزاری SPSS را بخواند، دستور خواندن به صورت ("مسیر و نام فایل") read.SPSS است. (موسوی ندوشن، ۱۳۹۱)

با اجرای دستور زیر محاسبات مربوط به مدل DINA انجام می شود از این مرحله به بعد دستورهایی برای مشاهده نتایج تحلیل باید اجرا شود.

```
dina<-dina(data=Datafile,qmatrix=Qmtx,rule="DINA")
```

برای محاسبه پارامترهای حدس و لغزش و خطاهای استاندارد مربوط به آن ها از دستور زیر استفاده می شود.

```
dina$coef
```

همچنین از دستور زیر می توان به منظور محاسبه احتمال چیرگی آزمودنی ها در هر یک از طبقات متغیر مکنون استفاده نمود.

```
dina$posterior
```

دستور زیر برای محاسبه احتمال وقوع هر یک از الگوهای متغیر مکنون و فراوانی مورد انتظار طبقات صفات استفاده می شود:

```
dina$attribure.patt
```

دستور

```
Print(dina)
```

الگوهای مهارتی که بیش ترین احتمال را دارند مشخص می نماید. و همچنین از دستور زیر برای محاسبه شاخص های اصلی برای هر یک از آیتم ها استفاده می شود (راوند^۱ و همکاران، ۲۰۱۳).

```
Summary(dina)
```

یافته ها :

در ادامه نتایج مدلسازی تشخیصی شناختی بر اساس مدل DINA با استفاده از داده های شبیه سازی شده ارائه شده است. برای انجام محاسبات پاسخهای ۱۰۰۰ نفر در یک آزمون ۳۰ سوالی شبیه سازی شده است. بر اساس ماتریس Q هر سوال به یک صفت مکنون مرتبط شده است. در مجموع ۵ صفت مکنون به عنوان صفات اصلی در این پژوهش در نظر گرفته شده اند. بخشی از ماتریس Q استفاده شده در این مثال

¹ Ravand

پژوهشی در شکل ۲ نشان داده شده است . در ادامه یافته های پژوهش بر اساس خروجی های نرم افزار R که در بخش روش آموزش داده شد آورده شده است.

جدول ۱ پارامترهای حدس و لغزش سوالات

سوال	حدس	لغزش	سوال	حدس	لغزش
سوال ۱	۰/۰۹۵	۰/۱۰۵	سوال ۱۶	۰/۱۲	۰/۰۹۹
سوال ۲	۰/۰۸۴	۰/۱۰۹	سوال ۱۷	۰/۰۸۱	۰/۱۲۴
سوال ۳	۰/۱۰۳	۰/۰۷۹	سوال ۱۸	۰/۰۹۴	۰/۰۴۹
سوال ۴	۰/۰۹۸	۰/۰۹۵	سوال ۱۹	۰/۰۸۴	۰/۰۶۸
سوال ۵	۰/۱۴۳	۰/۰۹۲	سوال ۲۰	۰/۱	۰/۱۲۲
سوال ۶	۰/۰۹۲	۰/۱۰۲	سوال ۲۱	۰/۱۰۷	۰/۱۰۲
سوال ۷	۰/۱۱۱	۰/۱۰۶	سوال ۲۲	۰/۱۰۶	۰/۱۴۸
سوال ۸	۰/۱۱۵	۰/۱۰۸	سوال ۲۳	۰/۱۰۷	۰/۰۶
سوال ۹	۰/۱۰۴	۰/۱۰۳	سوال ۲۴	۰/۱۰۶	۰/۱۰۸
سوال ۱۰	۰/۱۰۳	۰/۱۰۸	سوال ۲۵	۰/۰۹۶	۰/۰۶۲
سوال ۱۱	۰/۱	۰/۰۸۴	سوال ۲۶	۰/۰۸۵	۰/۰۶۳
سوال ۱۲	۰/۱۰۳	۰/۰۷۴	سوال ۲۷	۰/۰۸۹	۰/۱۰۹
سوال ۱۳	۰/۱۲۲	۰/۱۱۵	سوال ۲۸	۰/۱۰۹	۰/۱۳۲
سوال ۱۴	۰/۱۰۱	۰/۰۷	سوال ۲۹	۰/۰۸۶	۰/۱۴۵
سوال ۱۵	۰/۱۲۵	۰/۱۱۶	سوال ۳۰	۰/۱۰۵	۰/۱۶

جدول فوق پارامترهای حدس و لغزش براساس مدل DINA را نشان می دهد، براساس اطلاعات مندرج در این جدول پایین ترین ضریب حدس مربوط به سؤال های ۲ و ۱۹ با اندازه های ۰/۸۴ و بالاترین ضریب حدس مربوط به سوالهای ۵ و ۱۵ با مقادیر ۰/۱۴۳ و ۰/۱۲۵ می باشد. این ضرایب احتمال پاسخ گویی درست به سؤال برای دانش آموزانی است که به مهارت های مورد نیاز برای پاسخ گویی به سؤال تسلط ندارند.

همچنین پایین ترین مقدار لغزش مربوط به سؤالات ۲۳ و ۲۵ با مقادیر ۰/۰۶ و ۰/۰۶۲ است و بالاترین ضریب لغزش مربوط به سوالهای ۲۸ و ۲۹ با مقادیر ۰/۱۴۵ و ۰/۱۳۲ می باشد. این ضریب نشان دهنده احتمال پاسخ گویی غلط به سؤال برای دانش آموزانی است که به مهارت های مورد نیاز برای پاسخ گویی به سؤال تسلط دارند.

هرچه پارامترهای حدس و لغزشی در سؤال کوچک‌تر باشد، نشان‌دهنده برازش بهتر میان طرح سنجش تشخیصی، داده‌های تجربی و مدل‌شناختی است (راوند و همکاران ۲۰۱۳)

جدول ۲ ضرایب تشخیص و RMSEA سوالات

سوال	IDI	RMSEA	سوال	IDI	RMSEA
سوال ۱	۰/۸۰۰	۰/۰۶۹۹۶	سوال ۱۶	۰/۷۸۱	۰/۰۶۸۸۲
سوال ۲	۰/۸۰۷	۰/۰۶۳۳۵	سوال ۱۷	۰/۷۹۵	۰/۰۶۵۷۱
سوال ۳	۰/۸۱۸	۰/۰۶۳۴۳	سوال ۱۸	۰/۸۵۷	۰/۰۵۴۶۵
سوال ۴	۰/۸۰۸	۰/۰۶۵۹۷	سوال ۱۹	۰/۸۴۸	۰/۰۴۸۱۳
سوال ۵	۰/۷۶۴	۰/۰۶۳۲۵	سوال ۲۰	۰/۷۷۸	۰/۰۷۴۰۶
سوال ۶	۰/۸۰۶	۰/۰۷۸۴۶	سوال ۲۱	۰/۷۹۱	۰/۰۶۲۵۵
سوال ۷	۰/۷۸۳	۰/۰۷۶۱۵	سوال ۲۲	۰/۷۲۴	۰/۰۶۴۵۵
سوال ۸	۰/۷۷۷	۰/۰۶۹۰۶	سوال ۲۳	۰/۸۳۲	۰/۰۵۷۹۹
سوال ۹	۰/۷۹۳	۰/۰۶۴۷۷	سوال ۲۴	۰/۷۸۶	۰/۰۷۱۸۲
سوال ۱۰	۰/۷۸۸	۰/۰۶۴۶۶	سوال ۲۵	۰/۸۴۲	۰/۰۶۴۲۷
سوال ۱۱	۰/۸۱۵	۰/۰۵۳۱۸	سوال ۲۶	۰/۸۵۲	۰/۰۵۴۷۶
سوال ۱۲	۰/۸۲۳	۰/۰۶۵۸۶	سوال ۲۷	۰/۸۰۲	۰/۰۴۰۱۲
سوال ۱۳	۰/۷۶۳	۰/۰۷۹۹۱	سوال ۲۸	۰/۷۵۸	۰/۰۵۵۰۴
سوال ۱۴	۰/۸۳۹	۰/۰۶۱۴۲	سوال ۲۹	۰/۷۶۹	۰/۰۵۰۵۱
سوال ۱۵	۰/۷۵۹	۰/۰۶۹۸۲	سوال ۳۰	۰/۷۳۵	۰/۰۶۲۶۳

جدول شماره ۲ شاخص‌های تشخیص سؤال^۱ (IDI) و شاخص برازندگی سؤال (RMSEA) را ارائه می‌نماید. شاخص IDI رابطه معکوسی با شاخص‌های حدس و لغزش دارد و هرچه این شاخص کوچک‌تر باشند، شاخص بهتری تشخیص سؤال بالاتر است. دلاتوره و پارک (۲۰۱۲) نحوه محاسبه این شاخص برای هر سؤال را با استفاده از فرمول زیر بیان نموده‌اند.

$$IDI_j = 1 - S_j - g_j$$

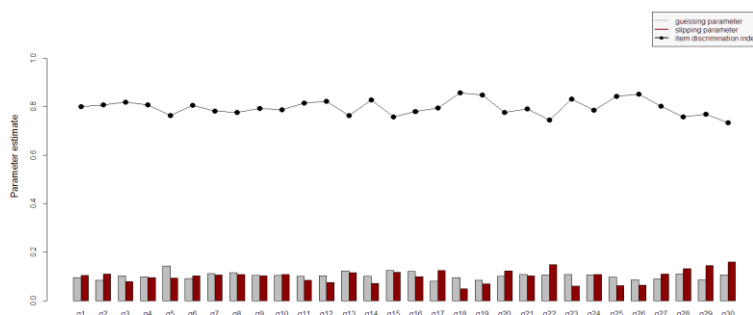
فرمول ۴

شاخص RMSEA نیز میزان برازندگی سؤال با مدل شناختی را نشان می‌دهد، این شاخص براساس توزیع

¹ Item Discrimination Index

حتی دولت هرچه کوچکتر باشد، نشان دهنده برارزش بالاتر سؤال با مدل است به صورت معمول سؤالات با شاخص RMSEA کوچکتر از ۰/۰۵ را سؤال دارای بهترین برازندگی با مدل و سؤالات با RMSEA بزرگتر از ۰/۱ را فاقد بر ارزش با مدل معرفی می نمایند.

براساس نتایج مندرج در جدول سؤالات ۱۸ و ۲۶ با شاخص های IDI برابر با ۰/۸۵۷ و ۰/۸۵۲ بهترین شاخص تشخیص و سؤالات ۲۷ و ۱۹ و با شاخص RMSEA برابر با ۰/۴۰۱۲ و ۰/۴۸۱۳ بیشترین میزان برارزش با مدل را دارا هستند



شکل ۳ پارامترهای حدس و لغزش و ضرایب تشخیص سؤالات

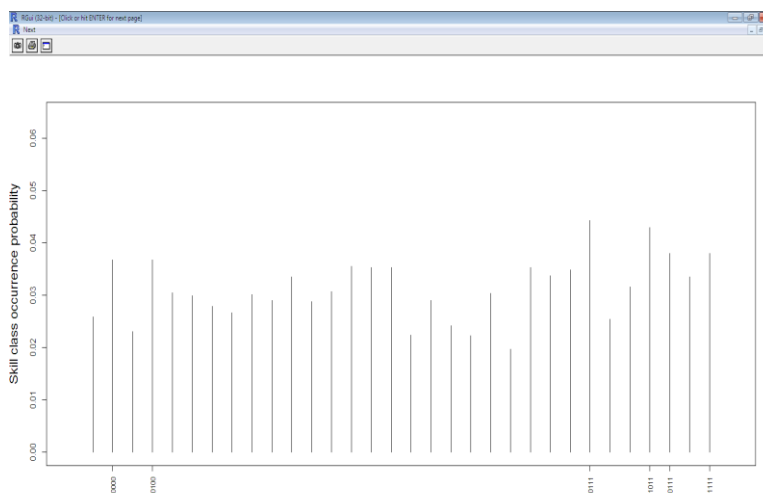
شکل ۳ پارامترهای حدس و لغزش و ضرایب تشخیص هر یک از ۳۰ سؤال را به صورت نمودار نمایش می دهد. همانگونه که در نمودار مشخص است هر چه حدس و لغزش کمتر باشند ، ضریب تشخیص سؤال بالاتر است.

جدول ۳ احتمال و فراوانی موردانتظار طبقات الگوی صفات

الگوی صفات	احتمال	فراوانی مورد انتظار	الگوی صفات	احتمال	فراوانی مورد انتظار
00000	۰/۰۲۵۷۷	۲۵/۷۷	11100	۰/۰۲۳۳۹	۲۳/۳۹
10000	۰/۰۳۶۶۸	۳۶/۶۸	11010	۰/۰۲۸۹۹	۲۸/۹۹
01000	۰/۰۲۲۹۷	۲۲/۹۷	11001	۰/۰۲۴۱۴	۲۴/۱۴
00100	۰/۰۳۶۷۵	۳۶/۷۵	10110	۰/۰۲۲۱۹	۲۲/۱۹
00010	۰/۰۳۰۴۵	۳۰/۴۵	10101	۰/۰۳۰۳۰	۳۰/۳۰
00001	۰/۰۲۹۸۷	۲۹/۸۷	10011	۰/۰۱۹۶۸	۱۹/۶۸
11000	۰/۰۳۷۸۷	۳۷/۸۷	01110	۰/۰۳۵۲۶	۳۵/۲۶
10100	۰/۰۲۶۶۱	۲۶/۶۱	01101	۰/۰۳۳۷۳	۳۳/۷۳
10010	۰/۰۳۰۱۰	۳۰/۰۹	01011	۰/۰۳۴۸۶	۳۴/۸۶

۴۴/۲۴	۰/۰۴۴۲۴	00111	۲۸/۹۳	۰/۰۲۸۹۳	10001
۲۵/۴۱	۰/۰۲۵۴۱	11110	۳۳/۴۸	۰/۰۳۳۴۸	01100
۳۱/۶۰	۰/۰۳۱۶۰	11101	۲۸/۷۰	۰/۰۲۸۷۰	01010
۴۲/۸۵	۰/۰۴۲۸۵	11011	۳۰/۶۷	۰/۰۳۰۶۷	01001
۳۷/۹۸	۰/۰۳۷۹۸	10111	۳۵/۴۹	۰/۰۳۵۴۹	00110
۳۳/۴۲	۰/۰۳۳۴۲	01111	۳۵/۳۰	۰/۰۳۵۳۰	00101
۳۷/۹۹	۰/۰۳۷۹۹	11111	۳۵/۲۸	۰/۰۳۵۲۸	00011

جدول شماره احتمال وقوع هریک از طبقات مهارت براساس ارتباط تعریف شده میان سؤالات و صفات مکنون در ماتریس Q را نشان می‌دهد همچنین در این جدول فراوانی مورد انتظار برای هر طبقه از الگوی مهارت نیز ارائه شده است، براساس اطلاعات مندرج در جدول شایع‌ترین الگوهای صفات در نمونه مورد پژوهش الگوی (00111) با احتمال وقوع ۰/۰۴۴۲۴ و فراوانی مورد انتظار ۴۴/۲۴ نفر و الگوی (10011) با احتمال وقوع ۰/۰۱۹۶۸ و فراوانی مورد انتظار ۶۸/کم‌ترین شیوع را در نمونه داشته‌اند. در مجموع ۰/۰۳۷۹۹ درصد احتمال تسلط به همه مهارت‌ها (۱ ۱ ۱ ۱) و ۰/۰۲۵۷۷ درصد احتمال عدم تسلط در همه مهارت‌ها (۰ ۰ ۰ ۰) وجود دارد.



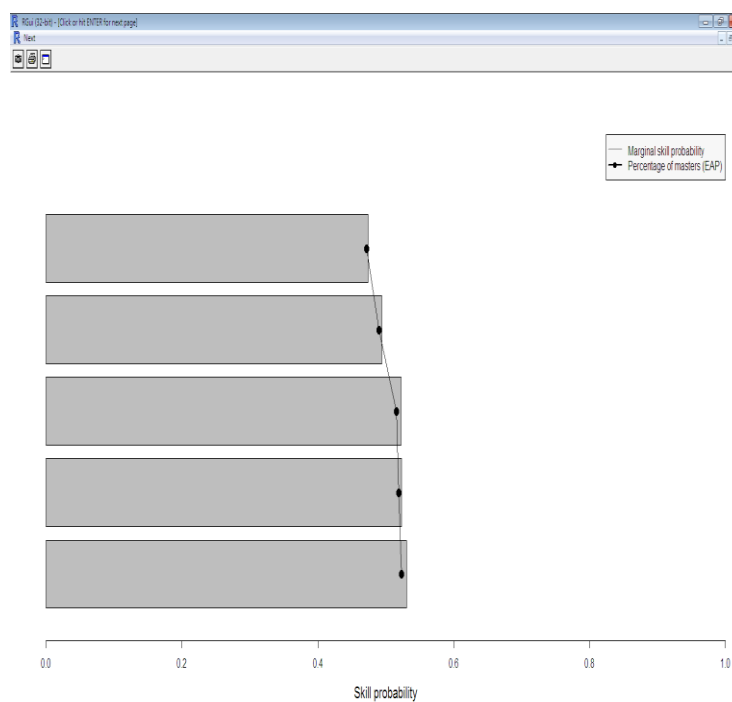
شکل ۴ احتمال کناری طبقات الگوی صفات

شکل ۴ احتمال هر یک از طبقات الگوی صفات را نمایش می‌دهد.

جدول ۴ احتمال کناری تسلط به صفات پنج گانه در نمونه مورد مطالعه

صفت	احتمال
صفت ۱	۰/۴۷۳
صفت ۲	۰/۴۹۴
صفت ۳	۰/۵۲۲
صفت ۴	۰/۵۲۲
صفت ۵	۰/۵۳۰

جدول شماره ۴ احتمال کناری تسلط به هر یک از مهارت های پنج گانه را نمایش می دهد با توجه به نتایج جدول صفت ۵ بیشترین احتمال با احتمال ۰/۵۳ و صفت ۱ کمترین احتمال با احتمال ۰/۴۳ دارا هستند، به عبارت دیگر ۵۳ درصد شرکت کنندگان در مهارت ۵ و ۴۷ درصد در مهارت یک به تسلط رسیده اند.



شکل ۵ احتمال کناری مهارت های پنج گانه

شکل شماره ۵ احتمال هر یک از صفات را به صورت نمودار نمایش می دهد

جدول 5 همبستگی تتراکوریک صفات پنج گانه

	صفت ۱	صفت ۲	صفت ۳	صفت ۴	صفت ۵
صفت ۱	۱	۰/۰۴۴	-۰/۰۸۱۱	-۰/۰۱۵۹	۰/۰۱۲۶
صفت ۲	۰/۰۴۴	۱	-۰/۰۳۹	۰/۰۵۶	۰/۰۴۳
صفت ۳	-۰/۰۸۱	-۰/۰۳۹	۱	-۰/۰۰۶	۰/۰۴۷
صفت ۴	-۰/۰۱۶	۰/۰۵۶	-۰/۰۰۶	۱	۰/۰۵۵
صفت ۵	۰/۰۱۳	۰/۰۴۳	۰/۰۴۶	۰/۰۵۵	۱

جدول شماره ۵ همبستگی تتراکوریک میان صفت‌ها را نشان می‌دهد همان‌گونه که در جدول مشخص است صفت‌های ۱ و ۳ دارای بالاترین همبستگی (منفی) با مقدار ۰/۰۸ و صفت‌های ۱ و ۵ دارای کم‌ترین همبستگی با مقدار ۰/۰۱۲ هستند.

بحث و نتیجه گیری

با استفاده از مدل‌های تشخیص شناختی (CDM) علاوه بر پروفایل‌های کاربردی برای آزمودنی‌ها، پارامترها و ویژگی‌های روان‌سنجی آزمون و سؤالات نیز برآورد می‌گردد. از آن‌جا که در پژوهش حاضر مدل DINA استفاده شده است پارامترها و ضرایب نیز که توسط این مدل برآورد می‌گردند در بخش یافته‌ها آورده شده‌اند دو پارامتر اصلی این مدل یعنی پارامترها حدس و لغزش میزان حدس‌پذیری سؤالات آزمون توسط آزمودنی‌هایی که به صفات مورد نیاز برای پاسخ‌گویی به آن سؤال تسلط ندارند (پارامتر حدس) هم‌چنین احتمال پاسخ‌گویی غلط به سؤال توسط آزمون‌هایی که به صفات مورد نیاز برای پاسخ‌گویی به آن سؤال تسلط ندارند (پارامتر لغزش) را مشخص می‌نمایند براساس یافته‌های این پژوهش کلیه سؤالات از امکان حدس و لغزش پائینی برخوردار بوده‌اند که با توجه به شبیه‌سازی بودن پاسخ‌ها طبیعی است.

هم‌چنین در این مدل یک ضریب تشخیص سؤال (IDI) برای هر سؤال ارائه می‌شود که نشان‌دهنده توانایی سؤال در تشخیص تسلط آزمودنی‌ها در صفات مکنون معرفی شده در ماتریس Q است در مثال حاضر کلیه سؤالات از ضریب تشخیص بالاتر از ۰/۷ برخوردار بوده‌اند.

یکی دیگر از ضرایب ارائه شده توسط این مدل ضریب RMSEA است که نشان‌دهنده میزان برازندگی سؤال با مدل شناختی تعریف شده است. هرچه این ضریب کوچک‌تر باشد برازندگی سؤال با مدل بیش‌تر است در مثال حاضر همه سؤالات دارای ضریب RMSEA کوچک‌تر از ۰/۰۵ بوده‌اند.

مدل DINA آزمودنی ها را در الگوهای صاف معرفتی شده براساس ماتریس Q طبقه بندی نموده و ضمن محاسبه احتمال هر کدام از الگوهای صفات فراوانی مورد انتظار هر الگو را نیز محاسبه می نماید. در این مثال کلیه طبقات از احتمال نسبتاً مشابهی برخوردار هستند. اما در پژوهش های تجربی حتماً نتایج طبقات با توجه به صفات مرتبط با هر الگو متفاوت خواهد بود.

یکی دیگر از خروجی های مدل DINA احتمال تسلط به هر یک از صفات (مهارت های) مکنون براساس نتایج حاصل از داده های تجربی است که استفاده از آن می تواند در طراحی دوره های آموزشی مکمل توسط برنامه ریزی این استفاده شود. در پژوهش حاضر بیشترین احتمال مربوط به صفت ۵ (۰/۵۳) و کمترین احتمال مربوط به صفت ۱ (۰/۴۷) است. با توجه به کاربردهای فراوان مدل های CDM در زمینه بررسی ویژگی های روان سنجی آزمون های قبلی، ساخت آزمون های براساس مدل های شناختی و غیرشناختی، تشخیص مشکلات یادگیری عمومی و اختصاصی در سطح مدرسه، کلاس و دانش آموز و ارائه بازخوردهای کاربردی به دانش آموزان، مدیران و سیاست گذاران استفاده از این مدل ها در علوم مختلف توصیه می شود. همچنین کاربرد مدل های مختلف CDM و مقایسه نتایج این مدل ها در جوامع یکسان و متفاوت و همچنین مقایسه با نتایج تئوری کلاسیک و تئوری سوال پاسخ می تواند زمینه مناسبی برای پژوهش های آتی باشد.

منابع

موسوی ندوشنی، سید سعید. (۱۳۹۱). آشنایی با زبان محاسباتی R. جزوه آموزشی. دانشگاه صنعت آب و برق (شهید عباسپور)

Aryadoust, V. (2011). Cogniittiive diiagnosttiic assessmenttt as an alltternattiive measurementt model. JALT Testing & Evaluation SIG Newsletter. March 2011. Vol. 15 #1. pp. 2-6.

Bartolucci, F. (2007). A class of multidimensional IRT models for testing unidimensionality and clustering items. *Psychometrika*, 72, 141-157.

Chen, J., & de la Torre, J. (2013). A general cognitive diagnosis model for expert-defined polytomous attributes. *Applied Psychological Measurement*, 37, 419-437.

De la Torre, J. (2009). DINA model and parameter estimation. *Journal of Educational and Behavioral Statistics*, 34, 115–130.

De la Torre, J. (2013). The generalized DINA model framework. *Psychometrika*.

de la Torre, J., & Douglas, J. (2004). Higher-order latent trait models for cognitive diagnosis. *Psychometrika*, 69(3), 333-353.

de la Torre, J., & Douglas, J. (2004). Higher-order latent trait models for cognitive diagnosis. *Psychometrika*, 69, 333–353.

DiBello, L, V. Roussos, L, A. Stout, W. (2007) . Review of Cognitively Diagnostic Assessment and a Summary of Psychometric Models. *Handbook of statistics*, 26

Gierl, M, J. Alves, C. (2010) Using pricipeld test design to develop and evaluate a diagnostic mathematics assessment in grade 3 and 6. Paper presented

at the annual meeting of the American educational research association .
Denever , Co, USA

Henson, R., Templin, J., & Willse, J. (2009). Defining a family of cognitive diagnosis models using log-linear models with latent variables. *Psychometrika*, 74(2), 191-210.

Junker, B. W., & Sijtsma, K. (2001). Cognitive assessment models with few assumptions, and connections with nonparametric item response theory. *Applied Psychological Measurement*, 25, 258-272.

Ravand, H. Barati, H & Widhiarso, W. (2013) Exploring Diagnostic Capacity of a High Stakes Reading Comprehension Test: A Pedagogical Demonstration. *Iranian Journal of Language Testing* Vol. 3, No. 1

Robitzsch, A. Kiefer, T. George, A. Uenlue, A. (2013). Package 'CDM'.
<https://sites.google.com/site/alexanderrobitzsch/software>

Rupp, A. A. Templin J, L. (2007). Review Article: Unique Characteristics of Cognitive Diagnosis Models. *Educational and Psychological Measurement*

Von Davier, M. (2005). A general diagnostic model applied to language testing data (RR-05-16). Princeton, NJ: Educational Testing Service.

Von Davier, M. (2008). A general diagnostic model applied to language testing data. *British Journal of Mathematical and Statistical Psychology*, 61, 287-307.

of attribute mastery in Massachusetts, Minnesota, and the U.S. National sample using the TIMSS 2007. *International journal of testing*, 11:2, 144-177.

Young-sun, L. & Yoon Soon, P. (2011) . A cognitive diagnostics modeling

