

رویکردهای نوین در مدیریت ورزشی

دوره ۵، شماره ۱۸، پاییز ۱۳۹۶

ص ص: ۸۹-۱۰۳

مقدماتی بر پژوهش‌های تجربی در مدیریت ورزش

علی محمد امیر تاش *

استاد گروه آموزشی تربیت بدنی و علوم ورزشی، دانشگاه آزاد اسلامی، واحد علوم و تحقیقات تهران، تهران، ایران

(تاریخ دریافت: ۱۳۹۵/۱۱/۰۹، تاریخ تصویب: ۱۳۹۶/۰۳/۱۰)

چکیده

هدف پژوهش حاضر جست‌وجو در مبانی نظری و عملی آزمون فرضیه، خطاهای تجربی و تصمیم‌گیری آماری، راه‌های کاهش این خطاها، تأکید بر تفاوت بین معناداری آماری با ارزشمندی یافته‌ها و اهمیت برآورد پارامترهای جامعه بود. نتایج نشان داد که استنباط‌های آماری ممکن است با خطاهای تصمیم‌گیری آماری نوع اول، آلفا، و نوع دوم، بتا، و اشتباهات تجربی همراه شوند. احتمال خطاهای نوع اول و دوم را پژوهشگر تعیین می‌کند، ولی اشتباهات تجربی، ریشه‌های انسانی، محیطی و روشی دارند و مقدار آن از طریق انحراف استاندارد معلوم می‌شود. از مقایسه عدد پی مربوط به آماره و عدد آلفا معناداری آماره معلوم می‌شود. ارزشمندی آماره از طریق مقدار واریانس‌های مشترک بین متغیرها ارزشیابی می‌شود. بروز خطای نوع اول جامعه را از اطلاعات مفید محروم می‌سازد، ولی بروز خطای نوع دوم اطلاعات غیرمفید را در ادبیات جاری می‌کند. اشتباهات تجربی شایان ملاحظه نیز، که می‌تواند کل اعتبار یافته‌ها را کاهش دهد، از طریق مقدار انحراف نشان داده می‌شود. ارائه ندادن برآورد پارامترهای جامعه نیز وضعیت متغیر مورد نظر را در جامعه مبهم می‌گذارد. پژوهش حاضر تلاشی است به منظور ارائه پیشنهادهایی برای کاهش اشتباهات در تفسیر و نتیجه‌گیری از پژوهش‌های تجربی، بالا بردن ارزشمندی یافته‌های تجربی و برآورد پارامترهای جامعه.

واژه‌های کلیدی

ارزشمندی، اشتباه، آماره، انحراف استاندارد، برآورد آماری، معناداری.

مقدمه

الگوی سنتی که در مجلات علمی-پژوهشی رشته ما در مدیریت تربیت بدنی و علوم ورزشی عمومیت و مقبولیت پیدا کرده است، جست‌وجو در روابط همبستگی، یا علت و معلولی، بین متغیرها و رسیدن به معناداری، یا غیرمعناداری آنهاست. نمونه‌های فراوانی از آنها را می‌توان در مجله‌های گرایش‌های مختلف رایج مشاهده کرد. مقاله حاضر از چنین سنتی تبعیت نمی‌کند، زیرا به اعتقاد نویسندگان که سال‌های زیادی در این وادی طلبه بوده است، چنین مقالات تجربی (آزمایشی) پیش از رسیدن به استنباط و نتیجه‌گیری زیرساخت‌های مهم‌تری نیاز دارند. در واقع، این مقاله قدمی است برای ارائه پاره‌ای از این زیرساخت‌ها تا چنین پژوهش‌هایی با اطمینان و اعتبار بیشتر انجام گیرند و نتایج قابل اعتماد تری به دست دهند. از این رو، الگوی سنتی یاد شده معیار مناسبی برای ارزشیابی تخصصی مقاله حاضر نخواهد بود، زیرا امیدواری به تأثیرات آن در افزایش کیفیت مطالعات تجربی، که تعداد آنها در مجلات علمی-پژوهشی ما رو به فزونی است، اولین هدف نویسندگان در نگارش آن بوده است. البته، آشنایی کافی با آمارهای استنباطی و آزمون فرضیه پیش‌نیاز استفاده بیشتر از مقاله حاضر است.

منابع زیادی برای تولید دانش در دانشگاه‌ها و مراکز پژوهشی در رشته مدیریت تربیت بدنی هزینه می‌شود. هدف پایانی همه این سرمایه‌گذاری‌های مادی و معنوی دسترسی به اطلاعات ناب است که اعتبار (۹، ۱۰) و ارزشمند بودن (۱۹، ۱) یافته‌های آنها قابل تأیید باشد.

روش‌های پژوهش و فنون آماری دو مجموعه ابزار علمی‌اند که تولید دانش نو را، در محدوده معینی از خطا، در مطالعات تجربی مدیریت تربیت بدنی ممکن می‌سازند. به علت ارتباط تنگاتنگ این دو با یکدیگر، اصولاً یکی بدون دیگری معنا و کاربرد پیدا نمی‌کند. به بیان دیگر، چنانچه

تعریف این نویسنده را از "روش‌های پژوهش" به شرح زیر بپذیریم: "فرایند پژوهش عبارت است از توصیف متغیرها، یا یافتن همبستگی بین متغیرها، یا جست‌وجوی روابط علت و معلولی بین متغیرها، یا ترکیبی از آنها براساس اصولی منظم و کنترل‌شده"، مشاهده می‌شود که اصل و ماهیت پژوهش با "متغیر" از یک سو و با یک یا ترکیبی از روش‌های آماری "توصیفی"، "همبستگی"، و "علت و معلولی" از سوی دیگر، ارتباط تنگاتنگ دارد (۳). در واقع، فرایند تولید دانش نو از طرح مسئله، یا سؤال، بر اساس اصول و ضوابط روش‌های پژوهش، آغاز می‌شود و با تجزیه و تحلیل‌های آماری و تفسیر و نتیجه‌گیری از آنها خاتمه پیدا می‌کند.

تصمیم‌گیری روی نتایج آماری و تعیین انواع و میزان خطاهایی که پژوهشگر در اعلام نتایج کار خود می‌پذیرد، مبحث ظریف و کم‌وبیش پیچیده‌ای است که تأثیرات اشتباه در آن در بعضی گزارش‌های پژوهشی و پایان‌نامه‌های تحصیلی قابل مشاهده است. بخشی از این پیچیدگی به رابطه معکوس بین بعضی معیارهای تصمیم‌گیری آماری مربوط می‌شود. برای نمونه، میزان احتمال بروز خطای تصمیم‌گیری آماری نوع اول (α -آلفا) با نوع دوم آن (β -بتا) رابطه معکوس دارد. همچنین، اندازه عدد بحرانی و اشتباهات تجربی (انحراف استاندارد)، هر دو، با افزایش تعداد آزمودنی‌ها کاهش می‌یابد، یا اینکه اندازه عدد مشاهده شده با اندازه احتمال وقوع آن (عدد "پی-P")، که در اصطلاح سطح معناداری نامیده می‌شود، مخالف جهت یکدیگرند. در ضمن مقایسه عدد مشاهده با عدد بحرانی برای رد کردن یا رد نکردن فرضیه صفر (پوچ) در آزمون‌های آماری یکسویه و دوسویه تفاوت دارد (۳). همه اینها نمونه‌هایی هستند که می‌توانند منابعی برای ایجاد تناقض ذهنی و بروز اشتباه نتیجه‌گیری از فعالیت‌های تجربی در بعضی پژوهشگران جوان‌تر شوند. از

مقدار کمبود اطمینان نتایج یافته‌های پژوهش‌های تجربی را با احتیاط بیشتری ارزشیابی کنند و از سوی دیگر، پژوهشگران آینده، با کاهش مقدار آنها، یافته‌های علمی دقیق‌تری را در اختیار مدیران تربیت بدنی قرار دهند.

برای کنترل یا کاهش اشتباه روی تعیین سطح معناداری نیز راهکارهایی وجود دارد. کم‌اطلاعی از معنا و مفهوم واقعی ریشه‌های این اشتباهات و راه کارهای موجود برای کنترل یا کاهش آنها، گاه به قیمت‌های گزافی برای جامعه تمام می‌شود. برای نمونه، نوعی عمل جراحی که می‌تواند با احتمال زیاد جان بیماران را نجات دهد، زیرا اشتباهاً از نظر آماری معنادار نبوده است توسط پژوهشگر بدون فایده اعلام می‌شود، یا نوعی درمان حیاتی بی‌فایده و بدون اثر که به‌علت انتخاب معیارهای نامناسب معنادار به دست آمده اشتباهاً تأیید می‌شود. بدیهی است تکرار هر دو این اشتباهات جامعه‌مصرف‌کننده را دچار ضرر و زیان فراوان می‌سازد (۲).

روش تحقیق

تحقیق حاضر از نوع "مروری" بوده و جزو تحقیقات توصیفی-تحلیلی است. روش جمع‌آوری اطلاعات کتابخانه‌ای و ابزار آن آثار مکتوب منتخبی از نویسندگان معتبر است که به عنوان مأخذ در روش‌های آماری و پژوهش شناخته می‌شوند (۷، ۱۱، ۱۲، ۱۳). به‌منظور دسترسی به اهداف مقاله حاضر بخش‌هایی از این منابع استفاده شده‌اند که به ماهیت آزمون فرضیه و اشتباهاتی که دقت آن را تهدید می‌کند مربوط می‌باشد. در همین راستا، راهکارهای ارائه شده که با کاهش این اشتباهات کیفیت یافته‌های علمی را تضمین می‌کند مورد تأکید قرار گرفته‌اند. در ادامه، نقش آماره در برآورد پارامتر جامعه مورد بررسی قرار گرفته تا هدف اصلی محاسبات آماری در روشن کردن چگونگی متغیر تحت بررسی در جامعه

اهداف مقاله حاضر کمک به کاهش پاره‌ای از همین ابهامات است.

نتیجه‌گیری‌های آماری، همانند هر نوع تحلیل و تفسیرهای ذهنی یا عملیاتی دیگر، فارغ از اشتباهات محیطی، انسانی، ابزاری، یا روشی نیستند. به همین گونه است که یکی از معانی "تجزیه و تحلیل آماری"، که به‌طور نمونه در روش‌های تحلیل واریانس و رگرسیون آماری انجام می‌گیرد، تجزیه واریانس کل به دو بخش سازنده آن (واریانس مربوط به اشتباهات و واریانس مربوط به ارتباط خالص بین متغیرها) است. شایان ذکر است که محور اصلی روش‌های آماری و پژوهش بر این اصل استوار است که بخش اشتباهات تا کمترین مقدار کاهش و بخش ارتباط خالص بین متغیرها تا بیشترین مقدار ممکن افزایش داده شود تا یافته‌ها با اعتبار بیشتر قابل استفاده شوند (۱۳). به‌طور کلی، اشتباهات از دو منبع اصلی سرچشمه می‌گیرند: ۱. خطاهای تصمیم‌گیری آماری (آلفا و بتا) و ۲. اشتباهات تجربی و عملی (انسانی، محیطی، ابزاری و روشی). اشتباهات تجربی و عملی که راه را برای ورود مجموعه‌ای از متغیرهای "مخل، اخلاک‌گر، رقیب، ناخواسته" به یافته‌های پژوهش باز می‌کنند، از طریق انحراف استاندارد (انحراف معیار) محاسبه می‌شوند. مقدار عددی این آماره هرچه بزرگ‌تر باشد، اشتباهات تجربی و عملی بیشتری در کار پژوهش وجود داشته است. برای کاهش این نوع اشتباهات راهکارهای آماری و روش‌های متعددی وجود دارد که می‌توان آنها را در بیشتر منابع روش پژوهش و آمار پیدا کرد (۳، ۴).

اشتباه تصمیم‌گیری روی تعیین سطح معناداری در تفسیر و نتیجه‌گیری از یافته‌های استنباطی بحث دیگری است که ریشه در احتمالات دارد. در واقع، احتمالات، مقدار درصد نبودن اطمینان به یافته‌های پژوهشی را نشان می‌دهد تا از یک سو، مصرف‌کنندگان با اطلاع از این

هدف تأمین شود. با این ترتیب، تحقیق حاضر رابطه همبستگی، یا علت و معلولی بین متغیرها را جست‌وجو نمی‌کند و در آن رابطه متغیر پیش‌بین با متغیر معیار، یا اثر متغیر مستقل بر متغیر تابع، مورد نظر نبوده است. به‌جای آن، هدف اصلی مقاله حاضر کاوش در ریشه‌های بروز اشتباه در مطالعات استنباطی و ارائه راهکارهایی برای مقابله با آنهاست تا چنانچه هدف محقق تجربی یافتن رابطه همبستگی یا علت و معلولی بین متغیرها می‌باشد، این اطلاعات را با دقت بیشتری پیدا کند. در نتیجه، یافته‌های علمی کاربردی‌تر شده و از اشاعه اطلاعات نادرست، یا بدون فایده، جلوگیری خواهد شد.

نتایج و یافته‌های تحقیق

فرضیه‌های آماری در برابر فرضیه‌های پژوهشی:

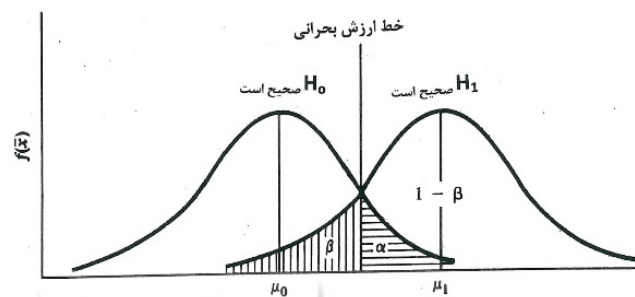
سؤالات و فرضیه‌های پژوهشی مطالعات تجربی حدس‌های نظری پژوهشگرند که در مقدمات گزارش به‌صورت مفهومی نوشته می‌شوند. این فرضیه‌ها به خودی خود قابل اندازه‌گیری و آزمایش نیستند و فقط در قالب فرضیه‌های آماری آزمایش می‌شوند. فرضیه‌های آماری که با حروف و علائم لاتین در کتاب‌ها نشان داده شده‌اند، به نام‌های متفاوت در ادبیات آمده‌اند. بعضی آنها را "فرضیه پوچ" و "فرضیه یک" و عده‌ای دیگر آنها را "فرضیه صفر" و "فرضیه مقابل" نامیده‌اند. شایان توجه است که این فرضیه‌های آماری در بسته‌های آماری کامپیوتری برای تعیین ردی و قبولی فرضیه پوچ جاسازی شده‌اند و نیازی به نوشتن آنها در گزارش‌های پژوهشی توسط محقق نیست. نکته مهم دیگر اینکه از بین دو نوع فرضیه آماری که نام برده شد، تنها فرضیه پوچ آزمایش می‌شود. از آنجا که این دو فرضیه در خلاف جهت یکدیگر قرار دارند، با تعیین شدن وضعیت فرضیه پوچ وضعیت ردی یا قبولی، فرضیه مقابل از آن استنباط می‌شود (۳، ۸، ۲۲).

منطق و فلسفه وجودی امور مادی جهان از دیدگاه آزمون فرضیه از دو حالت خارج نیست: یا وجود دارد یا وجود ندارد. کار آزمون فرضیه‌های آماری نیز تابع همین دو حالت است؛ یعنی فرضیه‌های آماری اثر، یا همبستگی، بین متغیرها یا تأیید می‌شود یا نمی‌شود. برای نمونه، می‌توان قبول کرد و به تجربه هم ثابت می‌شود، که بین قد و وزن رابطه وجود دارد، ولی عملکرد ورزشی با ساختار کف اقیانوس با احتمال زیاد بی‌ارتباط است، یا اینکه نوعی تمرینات ورزشی آمادگی جسمانی را تغییر می‌دهد یا تأثیری بر آن ندارد.

بر مبنای این منطق و واقعیت، اثر یا ارتباط همبستگی یا وجود دارد یا وجود ندارد. نبودن اثر یا ارتباط را در میحث آمار با فرضیه پوچ و بودن آن را با فرضیه یک، یا مقابل، نشان می‌دهند (۲۱). در شرایطی که فرضیه پوچ صحیح باشد، اثبات شدنی است که در تکرارهای خیلی زیاد، اندازه‌های اثر یا ارتباط همبستگی در نمونه‌های تصادفی هم تعداد دقیقاً مساوی و مشابه به‌دست نمی‌آیند. از حاصل این تفاوت‌ها توزیعی ساخته می‌شود که "توزیع فرضیه پوچ در شرایطی که این فرضیه آماری صحیح است" نام دارد. از آنجا که این توزیع از تفاوت‌هایی که همگی تصادفی بوده‌اند ساخته شده است، به آن توزیع احتمالات گفته می‌شود (۴، ۷، ۲۰). تمام اعداد بحرانی که ملاک رد یا قبول اثر یا ارتباط همبستگی قرار می‌گیرند، مانند t ، F و... از چنین توزیعی که به این صورت به وسیله خود آنها ساخته شده است، استخراج می‌شوند. این توزیع‌ها که به "توزیع نمونه‌های تصادفی هم تعداد" معروف‌اند، همه انواع نمونه با تعداد آزمودنی‌های متفاوت را که به نام "درجات آزادی" برای آزمون‌های یکسویه و دوسویه محاسبه و تبدیل می‌شوند، شامل می‌شوند (جدول‌های آنها را در پایان بیشتر کتاب‌های آمار می‌توان دید).

ارتباطها فقط دو نوع توزیع وجود دارد؛ یعنی مواردی که فرضیه پوچ صحیح است، و مواردی که فرضیه مقابل صحیح است. توزیع اول از نوع "توزیع احتمالات" هستند و به همین منظور برای تعیین احتمال وقوع یک شاخص آماری محاسبه شده، اختلاف یا همبستگی، توزیع احتمالات مربوطه برای آن به کار می‌رود (شکل ۱).

در شرایطی که فرضیه مقابل صحیح است (فرضیه پوچ غلط است)، بار دیگر از همه مواردی که اثر یا رابطه واقعی وجود دارد، برحسب مورد، توزیعی ساخته می‌شود که به آن "توزیع فرضیه مقابل در شرایطی که فرضیه مقابل صحیح است" می‌گویند. به این ترتیب، در طبیعت و به صورتی که خداوند آفریده است، برای همه اثرها و



شکل ۱- خطاهای نوع اول و دوم و توانایی آزمون در یافتن اثر یا ارتباط

همین صورت است و آماره "زد" فقط به عنوان نمونه آورده شده است (۶، ۱۶). از سوی دیگر، براساس تعریف می‌دانیم هر منحنی احتمالات از منهای بی‌نهایت تا مثبت بی‌نهایت ادامه (کشیدگی) دارد، یعنی دامنه‌های منحنی هیچ‌گاه محور ایکس‌ها را قطع نمی‌کنند (به صفر نمی‌رسد). به این صورت است که در توزیع احتمالات مربوط به صحیح بودن فرضیه پوچ، هیچ موردی، حتی اگر بسیار استثنایی (بزرگ) هم باشد، نمی‌تواند از آن خارج و به توزیع "فرضیه مقابل صحیح است" منتقل شود. معنای آن این است که در هیچ شرایطی هیچ‌گونه اثر یا همبستگی از نظر آماری پیدا نخواهد شد. برای حل این مشکل، پژوهشگر خطر پنج یا یک درصد اشتباه را می‌پذیرد و تصمیم می‌گیرد چنانچه احتمال وقوع این مورد استثنایی در توزیع "فرض پوچ صحیح است" فقط پنج یا یک درصد (برحسب مورد) بود، آن را به عنوان یک مورد نادر و کمیاب از این توزیع خارج و نتیجه‌گیری کند که فرضیه پوچ برای آن غلط است (اثر یا رابطه، برحسب مورد، وجود دارد). معنای "پذیرفتن خطر اشتباه به میزان پنج یا یک

می‌دانیم هر توزیع احتمالات که از تعداد بسیار زیاد آماره معین و فراوانی آنها به دست می‌آید، شکل طبیعی دارد و از این رو در این توزیع می‌توان بین درصدهای زیرمنحنی، که جمع آنها صددرصد یا یک است، با ارزش‌های آماره سازنده محور ایکس‌ها برای یافتن احتمال وقوع هر یک از این ارزش‌ها ارتباط برقرار کرد؛ یعنی با در دست داشتن یکی از آنها دیگری را به دست آورد. برای نمونه، چنانچه آماره مورد نظر عدد استاندارد "زد" که روی محور ایکس‌ها قرار دارد باشد، برای هر عدد "زد" می‌توان حجم زیرمنحنی را به دو مقدار احتمال بیشتر و کمتر از آن تقسیم کرد، یا اینکه مشخص کرد بین دو عدد "زد" چند درصد از حجم زیرمنحنی احتمالات قرار می‌گیرد (احتمال وقوع آنها چه مقدار است)، یا درصدهای بالاتر یا پایین‌تر از آنها چه مقدار هستند. برای مثال، می‌دانیم که ۶۸ درصد از فضای زیرمنحنی طبیعی بین مثبت و منفی یک انحراف استاندارد (معیار) از مقیاس "زد" قرار دارد، ۱۶ درصد از آن بیشتر و ۱۶ درصد هم از آن کمترند. برای سایر انواع آماره‌ها نیز وضعیت دقیقاً به

درصد" این است که اگر همین تجربه صد مرتبه تکرار شود، در پنج بار، یا یک بار، تصمیم پژوهشگر در تأیید اثر یا رابطه می‌تواند اشتباه باشد. بدین صورت تصمیمی که پس از آزمایش یک فرضیه آماری روی نتایج آن گرفته می‌شود، ممکن است به یکی از دو صورت زیر با اشتباه (خطا) همراه شود.

اشتباهات تصمیم‌گیری آماری:

الف) ممکن است H_0 (فرضیه پوچ) صحیح باشد، ولی تصمیم‌گیری محقق روی نتایج آزمون به اشتباه این باشد که باید آن را رد کرد (یعنی، $P \leq \alpha$ به دست آید). این را اشتباه نوع اول^۱ می‌نامند و با علامت α نشان می‌دهند.

ب) ممکن است H_0 غلط باشد، ولی تصمیم‌گیری محقق روی نتایج آزمون به اشتباه این باشد که باید آن را قبول کرد (یعنی، $P > \alpha$ به دست آید). این را اشتباه نوع دوم^۲ می‌نامند و با علامت β نشان می‌دهند.

توضیح مهم: تصمیم‌گیری و تفسیر نتایج باید فقط بر مبنای H_0 (فرضیه پوچ) صورت پذیرد، زیرا هیچ‌گونه معیار ریاضی یا آماری برای ارزشیابی فرضیه مقابل (H_1) وجود ندارد، از این رو این فرضیه اصولاً قابل آزمایش و تفسیر کردن نیست و نباید نتایج را بر مبنای آن تفسیر و نتیجه‌گیری کرد. البته فرضیه مقابل همیشه براساس نتیجه‌ای که از H_0 به دست می‌آید "استنباط" می‌شود. از دلایل مهمی که به روش‌های آزمون فرضیه "استنباطی" می‌گویند نیز همین استنباط فرضیه مقابل از روی فرضیه پوچ است. در ضمن، دلیل دیگر استنباطی بودن آزمون فرضیه نیز این است که اطلاعات جامعه از روی یافته‌های نمونه آن استنباط می‌شود (۱۲).

نقطه تعیین‌کننده احتمال α ، در توزیع احتمالات مواردی که H_0 صحیح بوده است^۳، نقطه‌ای از مرز اعداد

بحرانی را روی محور X ها نشان می‌دهد که چنانچه عدد مشاهده شده (تجربی) با آن مساوی یا بزرگ‌تر از آن شود، فرضیه پوچ (H_0) رد خواهد شد (رابطه یا اختلاف که وجود دارد از نظر آماری معنادار است). نقطه تعیین‌کننده احتمال β ، در توزیع احتمالات مواردی که H_0 غلط بوده است، نقطه‌ای از مرز اعداد بحرانی را روی محور X های این توزیع نشان می‌دهد که چنانچه عدد مشاهده شده (تجربی) کوچک‌تر از آن به دست آید، فرضیه پوچ (H_0) قبول خواهد شد (رابطه یا اختلاف وجود ندارد).

توضیح مهم دیگر اینکه P و α هر دو از همان جنس و مقیاس احتمالات هستند که مقدار اولی را کامپیوتر برای عدد محاسبه شده خود می‌دهد، و مقدار دومی را پژوهشگر برای تعیین سطح معناداری (عدد بحرانی را تعیین می‌کند) یافته‌هایش به عنوان معیار ردی، یا قبولی، فرضیه پوچ مشخص می‌سازد. در گزارش‌های پژوهشی همیشه باید علاوه بر عدد آلفا، که در مرحله طراحی توسط پژوهشگر به عنوان معیار رد یا قبول فرضیه پوچ تعیین می‌شود، عدد P مربوط به آماره محاسبه شده هم، که کامپیوتر در پایان محاسبات به دست می‌دهد، ارائه شود تا بتوان با مقایسه آنها نتیجه پژوهش را تعیین کرد (۱۷، ۱۶، ۸).

کاهش اشتباهات تصمیم‌گیری آماری:

واضح است که در آزمایش یک فرضیه آماری به دنبال تصمیمی هستیم که با کمترین اشتباه همراه باشد. شکل ۱ نشان می‌دهد که اشتباهات نوع اول و دوم از نظر مقدار با یکدیگر رابطه معکوس دارند. اینک، با توجه به رابطه معکوس بین دو نوع خطاهای تصمیم‌گیری یادشده، این سؤال مطرح می‌شود که تأکید روی کدامیک از این دو نوع خطا مناسب‌تر خواهد بود. به این صورت است که چنانچه تأکید را روی α بگذاریم (آنها را کوچک بگیریم)، مقدار β

1. Type One Error

2. Type Two Error

۳. به جزوه درسی "مقدمات آمار استنباطی" نگارنده مراجعه شود.

افزایش خواهد یافت. این موضوع نشان می‌دهد که ۱. هیچ یافته پژوهشی فارغ از خطای تصمیم‌گیری پژوهشگر نیست، و ۲. احتمال بروز یکی از انواع خطاهای تصمیم‌گیری را فقط به قیمت افزایش احتمال خطای نوع دیگر می‌توان کاهش داد (۱۷).

در آزمون آماری خیلی خوب α و β باید کوچک‌ترین مقدارهای ممکن را دارا باشند. البته با توجه به مطالب گفته‌شده، دست‌کم تا زمانی که حجم نمونه ثابت بماند، به چنین آزمونی نمی‌توان دست یافت. چون نتیجه کاربردی بروز خطای نوع اول مهم‌تر از خطای نوع دوم است، از این رو روش عملی مناسب به این ترتیب است که پژوهشگر میزان خطای نوع اول را از قبل در "سطح قابل اغماض" تعیین و سپس، سعی می‌کند برای چنین مقداری از خطای نوع اول آزمونی را انتخاب کند که در آن مقدار خطای نوع دوم کوچک‌ترین مقدار ممکن را داشته باشد (مثلاً استفاده از روش‌های آماری پارامتری به جای غیر پارامتری ...). تعیین "سطح قابل اغماض" برای احتمال خطای نوع اول، یک مسئله آماری نیست و مقدار آن با توجه به اهمیت کاربردی که رد کردن فرضیه پوچ دارد، معین می‌شود. معمولاً، سطح احتمالات کوچکی مانند ۰/۰۵ یا ۰/۰۱ را در انجام آزمون‌های آماری برای این نوع خطا به کار می‌گیرند. با توجه به تعبیری که از "احتمال" به دست می‌آوریم، اینکه مقدار خطای نوع اول برابر با ۰/۰۵ باشد، به این معناست که در هر ۱۰۰ آزمون مستقل، ۵ بار فرضیه پوچ را به غلط رد می‌کنیم (خطای تصمیم‌گیری نوع اول)؛ یعنی تأثیر یا ارتباطی را به اشتباه تأیید می‌کنیم (یافته‌ای را به کار می‌گیریم که در اصل غلط است).

توضیح اینکه نقطه‌های α و β هرکدام منحنی توزیع احتمالات مربوط به خود را به دو قسمت (یا سطح احتمالات) تقسیم می‌کنند. در توزیع احتمالات مربوط به

"فرضیه پوچ صحیح است" α ناحیه (احتمال) رد فرضیه پوچ و در توزیع احتمالات مربوط به "فرضیه مقابل صحیح است" β ناحیه (احتمال) رد کردن فرضیه مقابل را تعیین می‌کنند (شکل ۱). شایان توجه است که یکی از راه‌های کاهش همزمان α ، و به ویژه β ، در آزمون آماری افزایش حجم نمونه (N) است. واضح است که افزایش حجم نمونه با صرف هزینه بالا و وقت و انرژی زیادتر همراه است که مقرون به صرفه نیست، ولی راه مؤثری برای کاهش مقدار خطای نوع دوم به حساب می‌آید. در ضمن، معیار رد کردن H_0 مقایسه P که کامپیوتر از محاسبات خود درباره احتمال وقوع عدد مشاهده به دست می‌دهد و در توزیع احتمالات α قرار می‌گیرد و از همان جنس و مقیاس است، می‌باشد؛ یعنی اگر P مساوی یا کوچک‌تر از α تعیین شده در همان توزیع احتمالات "فرضیه پوچ صحیح است" به دست آمد، H_0 رد خواهد شد. شایان توجه است که P و α هرکدام، یک عدد بحرانی مشخص را روی محور X ها در توزیع احتمالات α (توزیعی که فرضیه پوچ در آن صحیح است) مشخص می‌کنند که α را پژوهشگر به عنوان معیار تصمیم‌گیری برای رد کردن و رد نکردن فرضیه پوچ تعیین می‌کند، ولی P را، که کامپیوتر به دست می‌دهد، مقدار اختلاف میانگین‌ها، یا بزرگی ارتباط همبستگی (عدد مشاهده)، مشخص می‌سازد. حال، این برداشت ممکن است ایجاد شود که هرچه P کوچک‌تر به دست آید، شدت تأثیر متغیر مستقل (فاصله میانگین گروه تجربی از میانگین گروه کنترل)، یا میزان ارتباط، بیشتر بوده است. مقدار α یا P فقط میزان احتمال خطای نوع اول را مشخص می‌سازد و ربطی به میزان تأثیر یا شدت متغیر مستقل یا میزان همبستگی ندارد (۱۳). میزان رابطه همبستگی، یا شدت تأثیر متغیر مستقل، از طریق شاخص‌های آماری دیگری مانند r یا ω^2 و ES و ... که "ارزشمندی" آماره را مشخص می‌سازند، قابل بررسی

هستند (۱). شایان توجه است که تنها در انواع پژوهش‌های علت و معلولی یا رابطه همبستگی فرضیه‌های آماری تدوین و به آزمایش گذارده می‌شوند.

کاهش اشتباهات تجربی:

منابع اصلی اشتباهات تجربی را متغیرهای اخلال‌گر (مخل، مزاحم، رقیب) از انواع انسانی (زمینه‌ای-ذاتی) یا محیطی تشکیل می‌دهند. نمونه‌هایی مانند انگیزش زیاد، خستگی، بی‌میلی، آمادگی (جسمی، روانی یا فنی)، تمرکز ذهنی، ملاحظات سیاسی-اجتماعی-اقتصادی، هوش و استعداد و... از انواع متغیرهای مخل انسانی غیرذاتی‌اند. دسته دیگر متغیرهای مخل انسانی مانند نژاد و جنسیت از نوع متغیرهای مخل انسانی ذاتی‌اند. در کنار آنها، متغیرهای مخل محیطی مانند ابزار اندازه‌گیری پیچیده یا غیرمعتبر، سروصدا، شلوغی یا درجه دمای محیط، زمان و مکان اندازه‌گیری، بی‌دقتی و عجله در جمع‌آوری اطلاعات خام اولیه، و... قرار دارد.

پژوهشگر متغیرهای اخلال‌گر را، به نسبت احاطه علمی که روی موضوع دارد، یا می‌شناسد یا از وجود آنها بی‌اطلاع است. در صورت شناخت، آنها را از طریق روش‌های مناسب آماری یا از طریق راهکارها و پیش‌بینی روش‌هایی که در زمینه روش‌های پژوهش و آمار وجود دارد، در حد امکان، کنترل می‌کند. در کارهای پژوهشی همیشه متغیرهای اخلال‌گری هم وجود دارند که در سطح علم روز، یا آگاهی پژوهشگر، ناشناخته باقی مانده‌اند. در این صورت تنها راه کنترل آنها استفاده از انواع روش‌های انتخاب تصادفی، که خود درجات متفاوتی از کیفیت تصادفی بودن را دارند، است (۵).

در صورتی که پژوهشگر متغیرهای مخل پژوهش خود را تشخیص دهد، برای کنترل بخشی از آنها روش‌های متعددی در علم آمار وجود دارد که تحلیل کوواریانس یک

نمونه شناخته‌شده آن است. می‌دانیم شاخص آماری این متغیرهای مخل همان انحراف استاندارد است که محاسبه آن در بخش آمارهای توصیفی کار پژوهشی بسیار متداولی است. این شاخص، علاوه بر آنکه میزان عددی اشتباهات تجربی، یا میزان خطای نمونه‌برداری را معین می‌کند، عملاً به عنوان "شاخص اشتباهات" (Error Term) با قرار گرفتن در مخرج بیشتر فرمول‌های آماری برای محاسبه میزان عددی آمارهای توصیفی و استنباطی، یا برآورد دامنه‌های تغییرات در پیش‌بینی‌های رگرسیونی یا پارامترهای جامعه از روی آمارها، نقش تعیین‌کننده دارد، که این نیز نوعی کنترل متغیرهای مخل به حساب می‌آید. افزایش تعداد آزمودنی و استفاده از روش‌های آماری پارامتریک به جای غیرپارامتریک راه‌های آماری مؤثر دیگری برای کنترل متغیرهای مخل هستند (۴).

در مبحث "روش‌های پژوهش" نیز راهکارهای دیگری برای کنترل بخش دیگری از متغیرهای مخل که پژوهشگر می‌شناسد وجود دارد. انتخاب گروه‌های مقایسه مشابه، همسان‌سازی گروه‌ها با استفاده از جفت‌های مشابه، استفاده از فقط یک گروه در اندازه‌گیری‌های متعدد (به کوواریانس توجه شود)، تقسیم گروه به زیرگروه‌های نزدیک به یکدیگر و... راه‌های دیگری برای کنترل این نوع متغیرهای مخل هستند.

آخرین دسته، متغیرهای مخلی‌اند که پژوهشگر می‌شناسد، ولی به علت تنگناهای اجرایی یا امکانات علم روز نمی‌تواند کاری برای کنترل آنها انجام دهد. در این صورت، به وجود و اثر احتمالی آنها در یافته‌های خود، به عنوان یک پژوهشگر مسئول، صادقانه اعتراف می‌کند. این اعترافات همان بخش "محدودیت‌های پژوهش" را که پژوهشگر معمولاً در فصل اول گزارش خود می‌نویسد، تشکیل می‌دهد.

همبستگی، یا اختلاف بین میانگین‌ها، باید خیلی بزرگ باشد تا برای مثال t محاسبه شده ارزش بزرگ‌تری از t بحرانی به‌دست آورد و از این‌رو فرضیهٔ پوچ رد شود.

- رد کردن H_0 علاوه بر α به N نیز بستگی دارد: هرچه N (درجهٔ آزادی را می‌سازد) بزرگ‌تر باشد، t بحرانی کوچک‌تر می‌شود و از این‌رو رد کردن H_0 آسان‌تر خواهد شد. به عبارت دیگر، اختلاف کمتری بین میانگین‌ها، یا مقدار همبستگی کمتری، لازم است تا H_0 رد شود. شایان توجه است که تصمیم دربارهٔ میزان احتمال اشتباه α بدون توجه به میزان احتمال اشتباه نوع دوم (β) گرفته می‌شود.

- عمل و اثر β بحث طولانی و مفصلی دارد که طرح همهٔ جزئیات آن در اینجا خارج از موضوع است. مختصر کاربردی از آن این است که:

- در یک آلفای معین ($\alpha=0.05$) مقدار عددی β را N (تعداد آزمودنی‌ها در گروه تجربی) و اختلاف واقعی بین 1μ و 2μ (میانگین‌های گروه‌های مورد مقایسه) معین می‌کنند.

- در یک α و N معین در گروه تجربی، مقدار احتمال β با افزایش اختلاف واقعی بین میانگین‌های گروه‌های کنترل و تجربی کاهش پیدا می‌کند (یعنی $1-\beta$)، که احتمال توانایی آزمون در پیدا کردن تأثیر متغیر مستقل، یا همبستگی بین متغیرها هنگامی که واقعاً وجود دارند، نام دارد، افزایش می‌یابد).

- توانایی آزمون ($1-\beta$) مقدار احتمال یک آزمون آماری در یافتن اثر یا همبستگی، هنگامی که آن اثر یا همبستگی واقعاً وجود دارد، را نشان می‌دهد.

به این احتمال در اصطلاح "توانایی آزمون آماری" (Power of the Test) گفته می‌شود. بعضی منابع به آن "توان آزمون"، که از نظر اینجانب مبهم است، هم گفته‌اند. برای افزایش آن توصیه می‌شود:

توجه: محدودیت‌ها، "قلمرو، مرز، یا محدودهٔ" پژوهش، که در بعضی ادبیات دیده می‌شود، نیستند، زیرا قلمرو، مرز و محدودهٔ مطالعه از طریق "اهداف ویژهٔ" آن باید تعیین شوند. این دو مقوله‌های بسیار متفاوتی‌اند.

اثر α ، β ، N ، اختلاف 1μ و 2μ ، و مقدار r در نتیجه‌گیری‌های

پژوهشی:

- α (اشتباه نوع اول): H_0 (فرضیهٔ پوچ) صحیح است، یعنی در اصل اختلافی بین میانگین‌ها، یا همبستگی بین متغیرها، وجود ندارد، ولی ما فرضیهٔ پوچ را با این مقدار احتمال اشتباه رد می‌کنیم. به عبارت دیگر، اختلاف یا ارتباطی را که در واقع وجود ندارد، اشتباهاً می‌پذیریم و تأیید می‌کنیم.

- β (اشتباه نوع دوم): H_0 (فرضیهٔ پوچ) غلط است، یعنی در اصل اختلاف میانگین‌ها، یا ارتباط بین متغیرها، حقیقی است، ولی ما فرضیهٔ پوچ را با این مقدار احتمال اشتباه قبول می‌کنیم. به عبارت دیگر، اختلاف، یا رابطه‌ای را که در واقع وجود دارد نمی‌پذیریم و تأیید نمی‌کنیم.

- با این ترتیب، یافته‌های علمی همراه با اشتباه تصمیم‌گیری آماری نوع اول نوعی قوانین علمی را توصیه می‌کنند که در اصل بی‌مورد، بی‌فایده و بی‌اثرند. در مقابل، یافته‌های علمی همراه با اشتباه تصمیم‌گیری نوع دوم انواعی از قوانین علمی را که در اصل می‌توانند برای رفع مسائل و مشکلات انسان مفید قابل استفاده و مؤثر باشند، غیرقابل توصیه اعلام می‌کنند.

- α را خود پژوهشگر تعیین می‌کند و ملاک آن اهمیتی است که کاربرد نتایج دارد (حیاتی یا مهم بودن آن). هرچه α کوچک‌تر شود، عدد بحرانی بزرگ‌تر می‌شود (مشاهدهٔ یک جدول اعداد بحرانی این را نشان خواهد داد) و از این‌رو رد کردن H_0 مشکل‌تر می‌شود؛ یعنی شدت

می‌یابد، زیرا با جمع‌تر شدن دامنه‌های منحنی توزیع گروه تجربی (به علت کاهش انحراف استاندارد که شاخص پراکندگی این گروه است) سهم کمتری از دامنه این توزیع، که احتمال "بتا" را تشکیل می‌دهد، وارد بخش مربوط به احتمال "آلفا" می‌شود و از این‌رو احتمال مربوط به "توانایی آزمون" افزایش پیدا می‌کند و اختلاف یا همبستگی مورد مطالعه از نظر آماری معنادار و تأیید می‌شود (شکل ۱).

- در هر مطالعه اختلاف بین میانگین‌ها، یا ارتباط همبستگی، در حد آلفای معین، میزان اشتباه نوع دوم (β) را فقط N تعیین می‌کند، یعنی هرچه N کوچک‌تر شود، ارزش β بزرگ‌تر می‌شود. از این‌رو احتمال ندیدن اختلاف یا ارتباط واقعی افزایش می‌یابد و برعکس.

معنادار بودن آماره در مقابل ارزشمند بودن آن:
معنادار بودن آماری یک شاخص محاسبه شده فقط گویای این است که با احتمال اشتباه آلفای ۱، یا ۵ درصد رابطه همبستگی یا تأثیر متغیر مستقل وجود دارد. این نتیجه‌گیری صرفاً آماری است و ارزش کاربردی این یافته‌ها را به دست نمی‌دهد. ارزش واقعی یک یافته آماری از نظر کاربردی مقدار واریانس خالص و بدون اشتباه متغیر تابع است که تحت تأثیر متغیر مستقل حاصل شده است (در مورد همبستگی نیز عبارت است از مقدار واریانس خالص و بدون خطا که متغیرهای تحت بررسی بر هم منطبق‌اند و در آن اشتراک دارند و به نام ضریب تعیین مشهور است). برای محاسبه هر یک از این واریانس‌ها روش‌های آماری ویژه وجود دارد (۱۰، ۱۹). برای نمونه، واریانس خالص ارتباط بین متغیرهای تحت بررسی در همبستگی را ضریب تعیین (توان دوم ضریب همبستگی) مشخص می‌سازد. بدین معنا که به اندازه ضریب تعیین محاسبه شده دو متغیر واقعاً با یکدیگر ارتباط دارند و یکی می‌تواند دیگری را به اندازه آن تعریف

۱. تعداد اعضای نمونه افزایش داده شود؛

۲. مقدار احتمال α افزایش داده شود؛

۳. شدت متغیر مستقل زیاد شود؛

۴. هر کجا که ممکن است، به جای آمارهای غیرپارامتری از انواع مشابه پارامتری استفاده شود (۱۵).

- اگر حکم (فرضیه مقابل) یکطرفه (یکسویه) به دوطرفه (دوسویه) تبدیل شود، α عملاً نصف می‌شود (کوچک‌تر می‌شود)، از این‌رو β بزرگ می‌شود و لذا احتمال ندیدن اختلاف یا رابطه واقعی زیاد می‌شود، زیرا $\beta-1$ کاهش می‌یابد (۱۴).

- هرچه اختلاف بین میانگین‌های تجربی و کنترل، یا دو گروه تجربی، زیادتر شود، احتمال رد کردن H_0 وقتی واقعاً غلط است، افزایش می‌یابد، یعنی β کوچک‌تر می‌شود و از این‌رو توانایی آزمون در تأیید اثر یا همبستگی واقعی افزایش پیدا می‌کند.

- در یک اختلاف معین بین میانگین‌ها و N معین، مقدار β با کاهش مقدار α افزایش خواهد یافت. بنابراین، در صورتی که α خیلی سخت‌گیرانه باشد (کوچک باشد)، احتمال ندیدن یک اختلاف یا ارتباط واقعی (β)، حتی اگر این اختلاف یا همبستگی به‌طور شایان ملاحظه‌ای هم بزرگ باشد، افزایش می‌یابد. دلیل آن این است که عدد بحرانی با کوچک شدن آلفا افزایش و (β) هم به همین دلیل افزایش پیدا می‌کند؛ از این‌رو رد کردن H_0 مشکل‌تر می‌شود و احتمال یافتن اثر متغیر مستقل یا همبستگی بین متغیرها که می‌تواند وجود داشته باشد، کم می‌شود.

- با افزایش N ، که موجب کاهش میزان اشتباهات تجربی (انحراف استاندارد)، که در همه محاسبه‌های آماری وجود دارد، می‌شود، توانایی آزمون ($\beta-1$) در رد کردن فرضیه پوچ غیرصحیح (احتمال پیدا کردن تأثیر یا رابطه هنگامی که واقعاً وجود دارد) افزایش

واریانس، مجذورکای (خی‌دو) و... اطلاعات ارزشمند دیگری هستند که بدون هزینه‌های اضافی از همان دسته اطلاعات خام قابل استخراج‌اند و کار پژوهشگر را کامل می‌کنند. بی‌تردید تنوع و وسعت اطلاعات پژوهشگر درباره ابزار و روش‌های آمار کاربردی امکان انتخاب بیشتری را، نه تنها در محاسبات آماری، بلکه در تدوین جامع‌تر اهداف ویژه مطالعه‌اش فراهم می‌آورد (۱۹).

آماره در مقابل پارامتر:

عامل مهم دیگری که در افزایش ارزش یافته‌های آماری نقش چشمگیری دارد، تعمیم آماره‌های حاصل از نمونه به پارامترهای جامعه مربوطه است. یک نمونه متداول و شناخته شده آن برآورد دامنه‌ای، یا نقطه‌ای در صورتی که اشتباه استاندارد میانگین صفر باشد، میانگین جامعه از روی میانگین نمونه است. در این محاسبه میانگین نمونه با علامت جبری $\pm$ با حاصل ضرب اشتباه استاندارد میانگین در عدد (Z) مربوط به $1/2$ آلفایی که پژوهشگر برای احتمال خطای آماری نوع اول خود انتخاب می‌کند (از جدول‌های تبدیل احتمالات به عدد Z استخراج می‌شود) جمع و تفریق می‌شود تا دامنه مورد نظر به دست آید؛ بدین معنا که چنانچه میانگین مورد نظر جامعه با آنچه در نمونه پیدا شده متفاوت است، با احتمال آلفای تعیین شده در این دامنه قرار خواهد گرفت. بدیهی است چنانچه اشتباهات صفر باشد، میانگین جامعه برابر با میانگین نمونه خواهد شد. البته چنین رویدادی بسیار نادر است و این مطلب فقط برای توضیح بیشتر بیان شد. این اطلاعات از این نظر بسیار ارزشمند است که پژوهشگر را به هدف نهایی آمار استنباطی، که تعمیم‌یافته از نمونه به جامعه هدف است، می‌رساند.

برآورد دامنه‌ای پارامترهای جامعه به میانگین ختم نمی‌شود. در واقع، برای هرگونه آماره‌ای که از محاسبه‌های

کند. در مقوله رابطه علت و معلولی، در طرح‌های پیش‌آزمون و پس‌آزمون، مقدار خالص اثر متغیر مستقل بر تابع از طریق محاسبه درصد افزایش، که نسبت میانگین افزایش بر میانگین پیش‌آزمون است، به دست می‌آید. در طرح‌های تجربی مقایسه دو گروه مستقل محاسبه مجذور امگا (ω^2)، که راه محاسباتی ویژه خود را دارد، تفاوت واقعی دو گروه را روی متغیر تابع مورد نظر به دست می‌دهد. به این معنا، که چنانچه در شرایط معین همان تجربه متغیر "الف" اتفاق بیفتد، متغیر "ب" به وجود خواهد آمد (۳).

مقدار مجذور امگا برای انواع تحلیل واریانس هم از طریق فرمول‌های ویژه خودشان قابل محاسبه‌اند و تحلیل و تفسیر آنها هم به همین صورت است. شایان توجه است در مواردی که شاخص آماری معنادار به دست می‌آید، محاسبه شاخص ارزشمندی (Meaningfulness) مربوط به آن ضرورت بیشتر پیدا می‌کند و اطلاعات تکمیلی کاربردی‌تری را به دست می‌دهد. این شاخص‌ها انواع متفاوت دارند. علاوه بر آنچه گفتیم، یکی دیگر از آنها محاسبه "اندازه اثر" (Effect Size-ES) است. اندازه اثر از تقسیم اختلاف بین میانگین‌های تجربی و کنترل تقسیم بر انحراف استاندارد گروه کنترل به مقیاس اعداد استاندارد به دست می‌آید. مقدار آن تا $0/2$ کم، تا $0/5$ متوسط و تا $0/8$ بزرگ است (۱۸).

ملاحظه می‌شود که با یافتن معناداری یا غیرمعناداری (چون هر دو به یک اندازه اطلاعات مفید به دست می‌دهند)، نه تنها کار پژوهشگر به پایان نمی‌رسد، بلکه فصل و مرحله جدیدی از اطلاع‌رسانی برای او آغاز می‌شود. برای نمونه، محاسبه مجذور اتا (η^2) که همبستگی بین یک متغیر اسمی و یک متغیر کمی (در انواع تحلیل واریانس) را به دست می‌دهد، یا آزمون‌های تعقیبی، برای یافتن ریشه‌های معنادار شدن در تحلیل

توصیفی، همبستگی یا روابط علت و معلولی به دست می‌آید، می‌توان به همین صورت برای پارامترهای جامعه برآوردهای دامنه‌ای به دست آورد. کافی است پژوهشگر در مورد آن آگاه باشد و محاسبه آن را از کامپیوتر خود بخواهد. برای تحلیل و تفسیر آنها هم مانند میانگین که شرح داده شد، عمل می‌شود. در اهمیت این برآورد همین بس که چنانچه نتیجه‌های پژوهش در جامعه کاربرد پیدا نکنند، زحمت و هزینه‌های بیهوده‌ای صرف آن شده است. شایان ذکر است که در نرم‌افزارهای "اس-پی-اس-اس" این برآوردهای دامنه‌ای برای هر شاخص آماری، معمولاً در انتهای جدول‌ها، برای آلفا پنج و یک درصد داده می‌شود (۲۲).

گزارش پژوهش:

نقطه پایانی هر پژوهش انتشار نتایج آن است. پایان‌نامه تحصیلی و مقاله انواع رایج تر و روزآمدتر انتشارات هستند که معمولاً منبع سایر پژوهشگران نیز قرار می‌گیرند. ترکیب و ساختارهای متفاوتی برای رساله و پایان‌نامه‌های تحصیلی مشاهده می‌شود که به طور متداول بدنه اصلی آنها دارای یک ساختار پنج فصلی است. هر یک از این پنج فصل یک هدف کلی و یک محتوای منطقی دارد که باید در طول تدوین پیوسته رعایت شود. در فصل اول که تحت عنوان "مقدمه و معرفی" نوشته می‌شود، پژوهشگر طرح کلی و اهداف کار خود را ارائه می‌دهد. فصل دوم، به نام "پیشینه پژوهش"، منتخبی از آنچه را که سایر پژوهشگران درباره پژوهش در دست اجرا در گذشته پیدا کرده‌اند، معرفی می‌کند. در فصل سوم، با نام "روش‌شناسی"، پژوهشگر معین می‌کند که طرح کلی و اهداف پژوهش خود را با استفاده از چه ابزار و روش‌هایی و در مورد کدام جامعه به اجرا درخواهد آورد. فصل چهارم، تحت عنوان "یافته‌های آماری"، نتایج اجرای پژوهش را

به صورت تجزیه و تحلیل‌های آماری شامل می‌شود. در فصل پنجم، که "بحث و نتیجه‌گیری" نام دارد، پژوهشگر یافته‌های آماری خود را به طور نظری تحلیل و تفسیر می‌کند و برای رسیدن به نظریه‌های علمی از آنها نتیجه‌گیری به عمل می‌آورد (۲).

هر فصل باید از یک ترکیب و محتوای منطقی پیروی کند. ترتیب منطقی فصل اول عبارت است از: مقدمه جامع برای معرفی پژوهش و جلب نظر خواننده، مبانی و زیربنای نظری فرضیه‌های پژوهشی، بیان مسئله، ضرورت و اهمیت استفاده از هر یک از یافته‌های پژوهش، اهداف و سؤالاتی که پژوهشگر به دنبال آنهاست، پیش‌فرض‌ها، محدودیت‌ها، و تعریف مفهومی و نظری واژه‌ها و اصطلاحات. فصل دوم از سه قسمت تشکیل می‌شود: ۱. توضیح هدف و روش گردآوری ادبیات پژوهشی در چند سطر برای هدایت فکر خواننده در حین مطالعه؛ ۲. بدنه اصلی ادبیات، و ۳. نتیجه‌گیری از پژوهش‌های ارائه شده. محتوای فصل سوم عبارت است از: معرفی محتوای فصل در چند سطر، شرح روش و طرح پژوهش، جامعه هدف، نمونه و روش نمونه‌گیری، تعریف عملیاتی متغیرها، ابزار اندازه‌گیری، اعتبار و پایایی ابزار اندازه‌گیری، روش جمع‌آوری اطلاعات خام، روش‌های آماری که به کار گرفته خواهند شد. فصل چهارم، پس از چند سطر معرفی محتوای فصل یافته‌های آماری را براساس شماره ردیف اهداف ویژه دربرمی‌گیرد. در فصل پنجم، پس از چند سطر معرفی محتوای فصل خلاصه‌ای از یافته‌های مهم پژوهش که پژوهشگر در نظر دارد آنها را تفسیر و نتیجه‌گیری کند، از نظر یادآوری، ارائه می‌شود. پس از آن در مبحث "بحث و نتیجه‌گیری" یافته‌ها، در چارچوب فرضیه‌ای پژوهشی مطالعه در دست اجرا، اهداف ویژه، مبانی نظری، ادبیات پیشینه و محدودیت‌های پژوهش، برای تولید نظریه‌های جدید یا تکمیل نظریه‌های موجود،

۲. وجود α در قوانین علمی موجب می‌شود تا جامعه از موارد یا اطلاعاتی استفاده کند که اصولاً در حد احتمال مربوطه مناسب یا صحیح نیستند، یعنی تأثیر روش، دارو، یا تمرینات ورزشی توصیه شده‌اند که عملاً بی‌اثر یا غیرقابل استفاده‌اند.

۳. وجود β در قوانین علمی موجب می‌شود تا جامعه از موارد یا اطلاعاتی که عملاً برای مفید یا قابل استفاده‌اند، در حد احتمال مربوطه محروم بماند، یعنی تأثیر روش، دارو یا تمرینات ورزشی که می‌توانسته‌اند برای مفید یا قابل استفاده باشند، رد شده و توصیه نشده‌اند.

۴. به این ترتیب، وجود α و β در یافته‌ها و قوانین علمی هرکدام جامعه را از یک سو متضرر می‌سازد؛ وقت، پول و انرژی که صرف یافتن آنها شده، عملاً کم بهره شده و دامنه رشد و توسعه علم در آن زمینه‌ها محدود می‌شود.

۵. با توجه به ارتباط معکوس بین α و β ، مشاهده می‌شود که هیچ‌گونه یافته علمی خالص و بدون اشتباه نیست و از این رو نتیجه‌گیری و کاربرد یافته‌های علمی باید با احتیاط و دوراندیشی صورت پذیرد.

۶. هنگامی یک کار پژوهشی به تکامل می‌رسد که علاوه بر بررسی‌های معناداری، ارزشمندی یافته، و دامنه‌ای که این یافته در جامعه هدف قرار می‌گیرد، ارائه شود.

۷. در نهایت اینکه، چنانچه یافته‌های علمی پس از تولید به صورت‌های مختلف مانند گزارش‌های علمی، سخنرانی، مقاله، و پوستر در اختیار جامعه مصرف کننده قرار نگیرد، سودی از آنها حاصل نخواهد شد.

مورد بحث و نتیجه‌گیری قرار می‌گیرند. در ادامه، پیشنهادهاى برخاسته از یافته‌های پژوهش، با استناد به هر یافته، و پیشنهاد برای پژوهش‌های بیشتر به منظور تکمیل اطلاعات موجود آورده خواهد شد.

نوع دوم انتشار نتایج پژوهشی که مورد نظر است، "مقاله" است. بدنه اصلی مقاله‌های مستخرج از رساله و پایان‌نامه تحصیلی معمولاً از سه بخش تشکیل می‌شود. در بخش اول، تحت عنوان "مقدمه" خلاصه‌ای از فصل‌های اول و دوم رساله نوشته می‌شود. بخش دوم، تحت عنوان "روش‌شناسی" فصل سوم را شامل می‌شود، و در نهایت در بخش سوم منتخب مورد نظر از فصل‌های چهارم و پنجم ادغام و ارائه می‌شود. مقاله‌هایی که از سایر انواع پژوهش تهیه می‌شود، کم و بیش از همین الگو پیروی می‌کنند (۲).

بحث و نتیجه‌گیری

۱. شایان توجه است که احتمال دیدن یک اختلاف واقعی (رد کردن H_0) به این مفهوم نیست که H_0 به‌طور مطلق صحیح یا غلط است. باوجوداین، بیشتر محققان، حتی برای نمونه‌های واقعاً کوچک، هنگامی که به یک حد معنادار بودن مورد نظر نمی‌رسند، تمایل دارند نتیجه‌گیری کنند که اختلاف یا ارتباط، وجود ندارد، یا شایان توجه نیست و از این رو فرضیه پوچ را صحیح قبول می‌کنند. با توجه به بحثی که درباره β گذشت، مشاهده می‌شود این نتیجه‌گیری ضمانت کافی ندارد. البته می‌توان یک ضابطه دومیاری را برای قبول H_0 در نظر گرفت: یعنی چنانچه α (احتمال اختلاف یا ارتباط) بیشتر از مثلاً ۰/۰۵ و β (احتمال عدم اتفاق اختلاف ارتباط همبستگی) کمتر از ۰/۰۵ بود، H_0 قبول شود.

منابع و مآخذ

۱. امیرتاش، علی محمد (۱۳۹۲). «مقدمات آمار استنباطی»، جزوه درسی، دانشگاه آزاد اسلامی، واحد علوم و تحقیقات تهران، ص ۳۰-۴.
۲. ----- (۱۳۹۴). «روش‌های پژوهش با تأکید بر مطالعات پیمایشی، میدانی، و مقاله‌نویسی»، تهران: انتشارات معاونت پژوهشی واحد علوم و تحقیقات، دانشگاه آزاد اسلامی تهران، ص ۱۳۳-۱۰۰.
۳. تامس، جری؛ نلسون، جک؛ سیلورمن، استیفن (۱۳۸۸). «روش تحقیق در فعالیت بدنی»، ترجمه‌علی محمد امیرتاش و فرشاد تجاری. ج ۱، تهران: چاپ واحد علوم و تحقیقات دانشگاه آزاد اسلامی تهران، ص ۶۵-۴۴.
۴. ----- (۱۳۹۰). «روش تحقیق در فعالیت بدنی»، ترجمه‌علی محمد امیرتاش و فرشاد تجاری. ج ۲، تهران: چاپ واحد علوم و تحقیقات دانشگاه آزاد اسلامی تهران، ص ۱۲۰-۱۰۶.
5. Anderson, B. F. (1971). "The psychology experiment, an introduction to the scientific method". 2nd Edition. Brooks/Cole Publishing Co., Belmont California, USA, pp: 15-32.
6. Best, J. W., Khan J. V. (2005). "Research in education". 10th Edition. Pearson International Publication, Los Angeles, USA, pp: 45-56.
7. Campbell, D. T., Stanley, J. C. (1963). "Experimental and quasi experimental designs for research". Rand McNally College Publishing Co., Chicago, USA, pp: 12-44.
8. Gaston, L. (2014). "Hypothesis testing made simple". 1st Edition. Leonard Gaston Publication, New York, USA, pp: 91-118.
9. Gay, L. R. (1999). "Educational research, competencies for analysis and application". 6th Edition. Charles E. Merrill Publishing Co. (A Bell and Howell Company), Columbus, pp: 116-145.
10. Gay, L. R., Geoffrey, M., P., W. A. (2008). "Competencies for analysis and applications". Pearson International Publications, New York, USA, pp: 67-93.
11. Issac, S., W., B. M. (1997). "Handbook in research and evaluation". 3rd Edition. EDITS Publishers, San Diego, USA, pp: 22-35.
12. Keppel, G. (2004). "Design and analysis, a researcher's handbook". Prentice-Hall, INC, New Jersey, USA, 108-160.
13. Kerlinger, F. N. (1999). "Foundations of behavioral research". Holt, Rinehart & Winston Inc. & Engage Learning, New York, USA, pp: 218-264.
14. Lehman, E. L., Romano, J. P. (2010). "Testing statistical hypotheses". 3rd Edition. Springer Publications, New York, USA, pp: 18-65.
15. McCall, R. B. (1975). "Fundamental statistics for psychology". 2nd Edition. Harcourt Brace Jovanovich, San Francisco, USA, pp: 16-44.
16. McCall, R. B. (2000). "Fundamental statistics for behavioral sciences". Harcourt Brace Jovanovich, San Francisco, USA, pp: 46-52.
17. Minium, E. W. (2010). "Statistical reasoning in behavioral sciences". 6th Edition. John Wiley & Sons Inc., USA, pp: 185-210.

18. Murphy, K. R., Myers, B., Wolach, A. (2014). "Statistical power analysis, a simple general model for traditional and modern hypothesis testing". 4th Edition. Taylor and Francis, UK, pp: 74-81.
19. Thomas, J. R. N., Jack K., Silverman, S. J. (2005). "Research methods in physical activity". 4th Edition. Human Kinetics, USA, pp: 30-48.
20. Udinsky, B., Flavian, O., Steven, J., Linch, S. W. (1981). "Evaluation resource handbook-gathering, analyzing, reporting data". EDITS Publishers, San Diego, USA, pp: 18-26.
21. Weakleim, D. L. (2016). "Hypothesis testing and model selection in the social sciences". 1st Edition. Guilford Press, New York, USA, pp: 72-78.
22. Wilcox, R. R. (2017). "Introduction to robust estimation and hypothesis testing". 4th Edition. Academic Press Publisher, Nikki Levy, San Diego, USA, pp: 92-104.

