



دوماهنامه علمی - پژوهشی

د ۹، ش ۱ (پیاپی ۴۳)، فروردین و اردیبهشت ۱۳۹۷، صص ۱۸۵-۲۱۱

جایگاه گروه اسمی سنگین در زبان فارسی: شواهدی از شم زبانی گویشوران بومی

حمیده معرفت^۱، مصطفی مهدی زاده^{۲*}

۱. استاد زبان‌شناسی کاربردی، دانشگاه تهران، تهران، ایران.

۲. دانشجوی دکتری زبان‌شناسی کاربردی، دانشگاه تهران، تهران، ایران.

پذیرش: ۹۶/۹/۹

دریافت: ۹۶/۶/۲

چکیده

پژوهش‌های پیشین در زمینه جایگاه گروه اسمی سنگین نشان می‌دهد که در زبان‌های «فعل - مفعول» همانند انگلیسی، سازه‌های بلندتر پس از سازه‌های کوتاه‌تر (پسایندسازی) ظاهر می‌شوند؛ ولی در زبان‌های «مفعول - فعل» همانند ژاپنی و کره‌ای، سازه‌های بلند قبل از سازه‌های کوتاه (پیشایندسازی) قرار می‌گیرند. در این گستره پژوهشی، بررسی زبان فارسی می‌تواند جالب باشد؛ زیرا این زبان دارای توالی کلمه‌ای نسبتاً آزاد است و از نظر موقعیت هسته در گروه‌ها، هم هسته - آغاز و هم هسته - پایان است. پژوهش‌هایی که این پدیده را در فارسی بررسی کرده‌اند شواهدی مبنی بر پیشایندسازی گروه اسمی سنگین در این زبان یافته‌اند؛ ولی غالب آن‌ها از داده‌های پیکره‌ای و عملکردی استفاده کرده‌اند که بیشتر نشان‌دهنده استفاده زبان در دنیای واقعی است. پژوهش حاضر علاوه بر بررسی داده‌های عملکردی، با استفاده از آزمون قضاوت دستوری، شم زبانی فارسی‌زبانان را نیز تحلیل کرده است. نتایج حاصل از یک آزمون تولیدی عملکردمحور و یک آزمون قضاوت دستوری - که به ۳۷ فارسی‌زبان داده شد - نشان می‌دهد با اینکه فارسی‌زبانان در آزمون قضاوت دستوری هم پیشایندسازی و هم پسایندسازی را قابل قبول می‌دانند، در آزمون تولیدی غالباً مفعول مستقیم را - که دارای نشانگر «را» است - بدون توجه به سنگینی آن قبل از مفعول غیرمستقیم قرار می‌دهند. این یافته‌ها از دو جنبه اهمیت دارند: ۱. استفاده هم‌زمان از داده‌های تجربی و شم زبانی در پژوهش‌های زبانی تصویر دقیق‌تری از پدیده مورد بررسی فراهم می‌کند، و ۲. در زبان فارسی میزان مشخص یا معین بودن مفعول مستقیم تناوب مفعول‌ها را تعیین می‌کند نه سنگینی یا سبکی آن‌ها.

واژه‌های کلیدی: تناوب کلمه‌ای، جابه‌جایی گروه اسمی سنگین، مفعول مستقیم/ غیرمستقیم، پیشایندسازی/ پسایندسازی، استفاده زبانی/ شم زبانی.

۱. مقدمه

هر زبانی دارای یک تناوب کلمه‌ای بی‌نشان^۱ است؛ مانند زبان انگلیسی که در آن توالی بی‌نشان «فاعل - فعل - مفعول» است. همچنین، درمورد برخی فعل‌های دومفعولی، مانند donate (اهدای کردن)، تنها توالی «مفعول مستقیم قبل از مفعول غیرمستقیم» دستوری تلقی می‌گردد (Mazurkewich, 1984; Pinker, 1989). با وجود این، دیده شده است که در مواردی که مفعول مستقیم بلندتر است مفعول غیرمستقیم از آن پیشی می‌گیرد؛ پدیده‌ای که به «اصل وزن - آخر^۲» شهرت یافته است؛ زیرا در زبان‌های ژرمنی، مانند آلمانی و انگلیسی، ساخت‌های بلندتر یا سنگین‌تر تمایل دارند بعد از ساخت‌های کوچک‌تر قرار گیرند (Arnold & et al., 2000; Ross, 1967; Stallings & et al., 1998; Wasaw, 1997). در ابتدا مدل‌های افزایشی تولید جمله^۳ (Garrett, 1980; Bock & Levelt, 1994) این پدیده را جهان‌شمول انگاشتند و ادعا کردند که ساخت‌های کوتاه از ساخت‌های بلند بیشتر در دسترس هستند و از این رو، قبل از ساخت‌های بلند قرار می‌گیرند؛ ولی با بررسی زبان‌های غیرژرمنی، به‌خصوص با استفاده از داده‌های پیکره‌ای، پژوهشگران مشاهده کردند که در زبان‌های «مفعول - فعل»، مانند ژاپنی (Hawkins, 1994; Yamashita & Chang, 2001) و کره‌ای (Choi, 2007)، عکس این اصل صادق است. به عبارت دیگر، این ساخت‌های بلندند که مورد پیشایندسازی قرار گرفته‌اند و قبل از ساخت‌های کوتاه قرار می‌گیرند. درنتیجه، توجه پژوهشگران از جهان‌شمول بودن این پدیده به سمت متفاوت بودن آن در زبان‌های گوناگون تغییر کرد و اصطلاح «جابه‌جایی گروه اسمی سنگین^۴» رایج شد.

در این گستره پژوهشی، نمود این پدیده در زبان فارسی می‌تواند از جهاتی بااهمیت باشد؛ زیرا این زبان خلاف زبان‌هایی مثل ژاپنی و کره‌ای همواره هسته - پایان نیست (در فارسی فقط گروه فعلی هسته است^۵) و همچنین خلاف زبان‌های هسته - آغاز مثل انگلیسی دارای تناوب کلمه‌ای نسبتاً آزاد است (Karimi, 2003). تاکنون علاوه بر پژوهش‌هایی که وزن دستوری را به‌صورت یک عامل نقشی (در کنار عوامل نقشی دیگر مانند معرفگی و جاننداری) در پدیده‌هایی

چون قلب نحوی^۱ بررسی کردند (راسخ مهند و قیاسوند، ۱۳۹۳). فقیری و سامولیان^۲ (۲۰۱۴) و فقیری و همکاران (۲۰۱۴) از معدود پژوهشگرهایی هستند که اختصاصاً به بررسی نقش وزن در جایگاه دو مفعول مستقیم و غیرمستقیم در زبان فارسی پرداختند. آن‌ها به این نتیجه رسیدند که تناوب مفعول مستقیم و غیرمستقیم در فارسی به میزان زیادی به نوع مفعول مستقیم بستگی دارد، نه محض بلندی یا سنگینی آن.

آنچه در پژوهش‌های گذشته به آن پرداخته نشده است «دانش زبانی» (Newmeyer, 2003: 682) یا شم زبانی گویشوران بومی زبان فارسی نسبت به این پدیده است. به عقیده علائی و همکاران (۱۳۹۶: ۳) نمی‌توان با تکیه صرف بر شواهد موردی و داده‌های پیکره‌ای نقش عوامل پردازشی و راندها قوه پردازشگر ذهنی را در چیدمان ساخت‌های نحوی بررسی و اثبات کرد. نیومیر (۲۰۰۳) نیز معتقد است که در پژوهش‌های زبان‌شناسی نباید تنها به داده‌های پیکره‌ای بسنده کرد؛ بلکه باید در کنار این داده‌ها از داده‌های قضاوتی^۳ گویشوران بومی که نشان‌دهنده شم زبانی آن‌هاست نیز استفاده کرد. در همین راستا و به منظور تکمیل یافته‌های پیشین و به دست آوردن تصویر کامل‌تری از تناوب مفعولی و همچنین پدیده جابه‌جایی گروه اسمی سنگین پژوهش پیش رو بر آن شد تا با تحلیل داده‌های حاصل از یک آزمون قضاوت دستوری در کنار یک آزمون تولیدمحور به بررسی این پدیده بپردازد. در حقیقت، هدف پژوهش حاضر پی بردن به این مسئله است که چیدمان مفعول‌های مستقیم (کوتاه و بلند) و غیرمستقیم در تولیدات زبانی گویشوران بومی فارسی چه تفاوتی با نحوه نمره دادن آن‌ها به همین جملات دارد، وقتی که از آن‌ها خواسته می‌شود با استفاده از شم زبانی خود آن‌ها را قضاوت کنند. به همین منظور فرضیه‌های ذیل مطرح شد:

۱. در آزمون تولیدمحور، گویشوران بومی فارسی جمله‌ها را صرف‌نظر از طول مفعول مستقیم و جایگاه مفعول‌ها در جمله‌های ارائه‌شده، با توالی «مفعول مستقیم - مفعول غیرمستقیم» به یاد می‌آورند.
۲. در آزمون قضاوت دستوری، گویشوران بومی فارسی هر دو توالی «مفعول مستقیم - مفعول غیرمستقیم» و «مفعول غیرمستقیم - مفعول مستقیم» را بدون توجه به طول مفعول مستقیم قابل قبول می‌انگارند.

۲. مروری بر ادبیات پژوهش

۲-۱. سنگینی^۱ یا وزن

بهاگل (1909)، به نقل از فقیری و سامولیان، (2014) اولین پژوهشگری بود که با ارائه شواهدی از اصلی به نام اصل وزن - آخر صحبت کرد. به گفته وی بر اساس این اصل در زبان‌های گوناگون سازه‌ها معمولاً از کوتاه به بلند در کنار هم قرار می‌گیرند. پس از بهاگل - که گفته بود وزن تنها به طول یا بلندی سازه اشاره دارد - تعریف‌های گوناگونی از مفهوم وزن یا سنگینی ارائه شد (Arnold & et al., 2000).

چامسکی^{۱۱} (1975) این عقیده را که وزن صرفاً به بلندی یا تعداد کلمه‌های یک سازه اشاره می‌کند به چالش کشید. او اظهار داشت که برای مثال جمله ۱ از جمله ۲ طبیعی‌تر به نظر می‌رسد (با اینکه گروه اسمی مفعول مستقیم در جمله ۱ بلندتر است)، و این گونه نتیجه گرفت که وزن یا سنگینی سازه‌ها فقط به تعداد کلمه‌های آن‌ها بر نمی‌گردد؛ بلکه از پیچیدگی آن‌ها در جمله سرچشمه می‌گیرد.

1. They brought all the leaders of the riot in.

2. They brought the man I saw in.

در حقیقت، تقریباً پنج دهه پیش‌تر، راس (1967) وجه تمایز مشابهی بین بلند بودن و پیچیده بودن قائل شده بود و گفته بود که هر ساخت بلندی لزوماً پیچیده نیست. برای نمونه از دو جمله زیر جمله ۴ به گفته بسیاری از انگلیسی‌زبانان قابل قبول نیست:

3. I called almost all of the men from Boston up.

4. *I called the man you met up. (p. 49)

به گفته راس گروه اسمی مفعول مستقیم در جمله ۴ دارای یک جمله‌واره^{۱۲} درونی^{۱۱} است و همین باعث می‌شود این گروه اسمی از گروه اسمی مفعول مستقیم در جمله ۳ پیچیده‌تر شود. واسو (1997) نیز با استفاده از داده‌های پیکره‌ای تعریف‌های گوناگونی از مفهوم سنگینی را بررسی کرد و این‌گونه نتیجه گرفت که همه این تعریف‌ها به یک میزان در پیش‌بینی تناوب کلمه‌ها موفق‌اند. به بیان دیگر، همان‌طور که آرنولد و همکاران (2000) اظهار داشتند، پدیده جابه‌جایی گروه اسمی سنگین آن‌قدر قوی است که تقریباً با هر نوع تعریفی از سنگینی سازگار است.

۲-۲. نقش جابه‌جایی گروه اسمی سنگین

پژوهشگران معتقدند که پیش‌سازسازی یا پس‌سازسازی گروه اسمی سنگین می‌تواند دو نقش داشته باشد؛ یکی اینکه به شنونده کمک می‌کند تا آنچه می‌شنود راحت‌تر پردازش کند، و دوم اینکه می‌تواند به گوینده کمک کند تا مفاهیم و آنچه را می‌خواهد بگوید بهتر و راحت‌تر برنامه‌ریزی کند. در میان پژوهش‌هایی که این موضوع را از منظر شنونده و پردازشگر بررسی کرده‌اند (Frazier & Fodor, 1978; Kimball, 1973) اصل شناسایی سریع سازه‌های بلافصل هاوکینز (1994) بسیار جالب است. بر اساس این اصل در شرایط برابر تناوبی ارجح است که طی آن پردازشگر بتواند همهٔ هسته‌های سازه‌های بلافصل فعل را سریع شناسایی کند. برای نمونه درمورد فعل‌های دومفعولی انگلیسی، پس‌سازسازی گروه اسمی سنگین و قرار دادن آن بعد از گروه حرف اضافه در جملهٔ ۶ باعث می‌شود پردازشگر هسته‌های این دو سازه را بعد از پردازش تنها چهار کلمه شناسایی کند (در حالی که این کار در جملهٔ ۵ با پردازش هشت کلمه امکان‌پذیر است)؛ ولی در زبانی مانند ژاپنی عکس این کار، یعنی پیش‌سازسازی گروه اسمی سنگین، باعث راحتی کار پردازشگر می‌شود (Chang, 2009).

5. I [VP introduced¹ [NP some² friends³ that⁴ John⁵ had⁶ brought⁷] [PP to⁸ Mary.]]

6. I [VP introduced¹ [PP to² Mary³] [NP some⁴ friends that John had brought.]]

با وجود این، در زبان فارسی این موضوع صدق نمی‌کند. نمونهٔ زیر از فقیری و سامولیان (2014: 228) نشان می‌دهد که پردازشگر جهت شناسایی سازه‌های وابسته به فعل در هر دو تناوب با تعداد کلمات یکسانی مواجه می‌شود:

۷. یوسف [VP یک^۱ کتاب^۲ آموزش^۳ عکاسی^۴] [PP از^۵ کتابخانه^۶] گرفت^۷.

۸. یوسف [VP از^۱ کتابخانه^۲] [NP یک^۳ کتاب^۴ آموزش^۵ عکاسی^۶] گرفت^۷.

همچنین، پس‌سازسازی گروه اسمی سنگین در زبانی مانند انگلیسی می‌تواند در مراحل برنامه‌ریزی و تولید زبانی کمک کند. صحبت کردن نیازمند این است که صور و ساختارهای متفاوت از حافظهٔ بازیابی شوند و بر اساس قوانین هر زبانی کنار هم قرار گیرند و تمام این مراحل باید با سرعت هر چه بیشتر انجام شود (Levelt, 1989). برای نمونه، آرنولد و همکاران (2000) جملهٔ زیر را که گزارشگران بیسبال بسیار به‌کار می‌برند نقل می‌کنند و اظهار می‌دارند که پس‌سازسازی گروه اسمی در تولید چنین جمله‌ای زمان بیشتری در اختیار گوینده قرار می‌دهد تا هنگام تعویض، نام بازیکن ورودی را در لیست بازیکنان پیدا کند و جملهٔ خود را بسازد.

9. That brings to the plate Barry Bonds.

علاوه بر این، مدل‌هایی که تولید جمله را از دیدگاه در دسترس بودن^{۱۲} بررسی می‌کنند، این‌طور می‌انگارند که سامانه تولید زبانی ما جمله‌ها را به صورت خطی و افزایشی آرایش می‌دهد و توالی سازه‌ها هنگام تولید به میزان در دسترس بودن آن‌ها بستگی دارد (Bock, 1982). پسایندسازی گروه اسمی سنگین در انگلیسی با این واقعیت همخوانی دارد؛ زیرا سازه‌های کوتاه‌تر راحت‌تر در دسترس قرار می‌گیرند و در نتیجه در هنگام تولید سریع‌تر بازیابی می‌شوند.

۳-۲. جهان‌شمولی پسایندسازی گروه اسمی سنگین

از آنجا که مدل‌های تولید زبانی مبتنی بر در دسترس بودن ریشه در اصول کلی قوه شناخت انسان دارند، می‌توان نتیجه گرفت که پسایندسازی یا پدیده «کوتاه قبل از بلند» یک پدیده جهانی است، ولی شواهد مغایر این نتیجه‌گیری هستند. یاماشیتا و چنگ (2001) نشان دادند در زبان ژاپنی پدیده پیشایندسازی رخ می‌دهد؛ یعنی افراد اغلب سازه‌های بلندتر را قبل از سازه‌های کوتاه‌تر قرار می‌دهند. آن‌ها با این باور که مدل‌هایی که پدیده پیشایندسازی یا پسایندسازی را صرفاً از منظر صوری توضیح می‌دهند به همه زبان‌ها قابل تعمیم نیستند از نظریه کدگذاری دستوری^{۱۳} (Bock & Levelt, 1994; Garrett, 1980) استفاده کردند و نظریه‌ای پیشنهاد دادند که هم بُعد صوری و هم بُعد مفهومی را دربر می‌گیرد. بر اساس این نظریه تصمیم‌های مربوط به تناوب سازه‌ها در مراحل گوناگون تولید جمله متفاوت است. به بیان دیگر، در مرحله‌ای که مفاهیم باید انتخاب شوند بُعد مفهومی و در مرحله‌ای که صور باید انتخاب شوند بُعد صوری بر این تصمیم‌ها اثر می‌گذارد. در بُعد مفهومی، سازه‌هایی که از لحاظ معنایی غنی‌تر و روشن‌ترند از سازه‌های دیگر پیشی می‌گیرند؛ در حالی که در بُعد صوری هر سازه‌ای که کوتاه‌تر و در دسترس‌تر باشد از دیگر سازه‌ها پیشی می‌گیرد. این دو بُعد که نمایانگر تمایلات متفاوتی هستند در زبان‌های گوناگون با هم رقابت می‌کنند و برتری هر کدام از آن‌ها تعیین‌کننده پیشایندسازی یا پسایندسازی است.

ژاپنی، برای نمونه، زبانی است ضمیرانداز و هسته - پایان که تناوب کلمه‌ای نسبتاً آزادی نیز دارد (خلاف انگلیسی) و در آن فعل بعد از موضوع‌هایش ظاهر می‌شود (مانند همه زبان‌های

فاعل - مفعول - فعل). در نتیجه از آنجا که سازه‌های گوناگون قبل از فعل قرار می‌گیرند، بُعد صوری یعنی فاصله بین فعل و هسته سایر سازه‌ها (Hawkins, 1994) اهمیت کمی در تناوب کلمه‌ای دارد. از این رو، در این زبان سازه‌هایی که بلندترند (و در نتیجه از لحاظ معنایی غنی‌ترند) پیشایندسازی می‌شوند (Chang, 2009). به بیان دیگر، زبان‌های ضمیرانداز با تناوب کلمه‌ای آزاد همانند ژاپنی، کره‌ای، و فارسی در تناوب کلمه‌ها بیشتر تحت تأثیر ویژگی‌های مفهومی سازه‌ها هستند.

۲-۴. جابه‌جایی گروه اسمی سنگین در زبان فارسی

یکی از کامل‌ترین پژوهش‌ها در مورد نقش وزن در جایگاه سازه‌ها در زبان فارسی به کوشش راسخ مهند و قیاسوند (۱۳۹۳) صورت گرفته است. آن‌ها با بررسی وزن سازه‌های گوناگون در حدود ۱۹۰۰ جمله استخراج‌شده از یک پیکره فارسی یافتند که سازه‌هایی که در جایگاه نهایی جمله قرار می‌گیرند به‌طور میانگین دارای وزن نسبی بیشتری هستند. آن‌ها همچنین در گام بعدی به تبیین نقش وزن سازه‌ها در پدیده قلب نحوی در این جمله‌ها پرداختند و مشاهده کردند که سازه‌هایی که مورد قلب نحوی قرار گرفته بودند (در ابتدای جمله و قبل از سازه فاعل ظاهر شده بودند) نسبت به سازه‌های مجاور در جمله که از روی آن‌ها حرکت انجام شده بود سنگین‌تر بودند. در گام بعدی آن‌ها به بررسی وزن سازه‌هایی که پسایندسازی شده بودند (یعنی به موقعیت بعد از فعل و انتهای جمله حرکت کرده بودند) پرداختند و نتیجه گرفتند که این سازه‌ها نسبت به سازه‌های مجاور که از روی آن‌ها حرکت صورت گرفته بود سنگین‌تر بودند. راسخ مهند و قیاسوند (۱۳۹۲) نیز با تحلیل یک پیکره گفتاری تلاش کردند تا به نقش عوامل کلامی گوناگون مانند وزن، جاننداری، و ساخت اطلاعی در پدیده پسایندسازی پی ببرند. در مورد نقش وزن در این پدیده یافته‌های آن‌ها نشان داد که هر چه سازه سنگین‌تر باشد احتمال پسایندسازی آن بیشتر است. نوع دیگری از جابه‌جایی که در زبان فارسی مورد بررسی قرار گرفته پدیده خروج بند موصولی است (شیخ‌الاسلامی، ۱۳۸۷؛ راسخ مهند و همکاران، ۱۳۹۱). راسخ مهند و همکاران (۱۳۹۱: ۱۶)، خلاف شیخ‌الاسلامی (۱۳۸۷) پس از بررسی داده‌های پیکره‌ای خود به این نتیجه رسیدند که هر چقدر نسبت طول بند موصولی بر گروه فعلی بیشتر باشد احتمال خروج و جابه‌جایی آن به جایگاه بعد از فعل بیشتر است.

اما بعد از چندین پژوهش پیکره‌ای - که در بالا به آن‌ها اشاره شد - علای و همکاران (۱۳۹۶) در رویکردی نوین به صورت تجربی و با استفاده از یک تکلیف بر خط خوانش جمله به بررسی نقش وزن دستوری سازه‌ها در دو پدیده جابه‌جایی بند موصولی و قلب نحوی کوتاه (از نوع پسایندسازی) پرداختند. آن‌ها با بررسی زمان خوانش جمله‌هایی که دارای سازه مورد نظر با سه سطح وزنی سبک، متوسط، و سنگین بودند پی بردند که جمله‌هایی که دارای سازه سنگین‌تر در جایگاه بعد از فعل بودند نسبت به جمله‌هایی که در آن‌ها جابه‌جایی صورت نگرفته بود دارای زمان خوانش کمتری بودند. به بیان دیگر، جابه‌جایی سازه‌های سنگین‌تر و قرار گرفتن آن‌ها در انتهای جمله باعث سهولت در پردازش این جمله‌ها و در نتیجه کمتر شدن زمان خوانش آن‌ها می‌شد.

۵-۲. تناوب مفعول‌ها در زبان فارسی

در زبان فارسی تا مدت‌ها تناوب مفعول‌ها بر اساس معیار مفعول‌نمایی افتراقی^{۱۴} (Karimi, 2003؛ راسخ مهند، ۱۳۸۳؛ مویدی و لطفی، ۱۳۹۲؛ میردهقان و یوسفی، ۱۳۹۵) توجیه می‌شد. این معیار در توالی‌های بی‌نشان تنها دو مکان برای مفعول مستقیم در نظر می‌گیرد: مفعول مستقیم مشخص با نشانگر «را» که قبل از مفعول غیرمستقیم و دور از فعل قرار می‌گیرد (جمله ۱۰)، و مفعول مستقیم نامشخص غیرصریح که بعد از مفعول غیرمستقیم و در کنار فعل قرار می‌گیرد (جمله ۱۱).

۱۰. علی کتاب را به او داد.

۱۱. علی به او کتاب داد.

فقیری و همکارانش با انجام پژوهش‌هایی مبتنی بر داده‌های پیکره‌ای و تجربی جایگاه‌های ممکن برای مفعول مستقیم را در فارسی دقیق‌تر توصیف کردند. فقیری و سامولیان (2014) پیکره‌ای را که از روزنامه همشهری استخراج شده بود بررسی کردند و بر اساس یافته‌های خود دیدگاه دوقطبی معیار مفعول‌نمایی افتراقی کریمی (2003) را ناقص خواندند؛ سپس تحلیلی پیوستاری از جایگاه مفعول مستقیم در فارسی ارائه دادند که در آن حداقل سه جایگاه (و به‌طور کامل چهار جایگاه) را برای مفعول مستقیم برشمردند. در یک سر این پیوستار که بر اساس میزان معین بودن گروه اسمی است مفعول مستقیم مشخص همراه با نشانگر «را» قرار

دارد، و در سر دیگر آن مفعول مستقیم غیرصریح، مفعول مستقیم نکره و مفعول غیرصریح وابسته نیز در میان این پیوستار قرار می‌گیرد.

فقیری و همکاران (۲۰۱۴) سپس از طریق یک پرسش‌نامه اینترنتی سعی بر جمع‌آوری داده‌های تجربی کردند تا نتایج به‌دست‌آمده از پژوهش پیکره‌ای خود را آزمایش کنند و همچنین نقش طول مفعول مستقیم در جایگاه ترجیحی آن را بررسی کنند. نتایج به‌دست‌آمده از این بررسی نتایج پژوهش قبلی را تأیید کرد و نشان داد که بلندی سازه همبستگی معنی‌داری با پیش‌سازای آن دارد. در مجموع، می‌توان این دو پژوهش کلیدی را در حوزه جایگاه مفعول مستقیم و پیش‌سازای گروه اسمی سنگین در فارسی این‌گونه خلاصه کرد که در این زبان میزان مشخص یا معین بودن مفعول مستقیم نسبت به بلندی و صرف وزن آن نقش پررنگ‌تری در تعیین جایگاه آن در جمله دارد. از این منظر، زبان فارسی بسیار شبیه زبان‌هایی مثل ژاپنی و کره‌ای است؛ زیرا در این زبان‌ها بُعد مفهومی نقش مهم‌تری از بُعد صوری دارد که نتیجه آن پیش‌سازای سازه سنگین است.

۶-۲. داده‌های پیکره‌ای در مقابل داده‌های قضاوتی در پژوهش‌های زبان‌شناسی

در زبان‌شناسی نظری دو مکتب رایج است: یکی زبان‌شناسی صوری - زایشی^{۱۵} که بر دانش زبانی (Manning, 2003)، شم زبانی، دوقطبی بودن^{۱۶}، و بررسی جمله‌های جدا از بافت تأکید دارد؛ و دیگری زبان‌شناسی نقش‌گرا^{۱۷} که بر استفاده زبانی، تأثیر بافت بر نوع دستور و جمله، و درجه‌بندی^{۱۸} تأکید می‌کند. این دو مکتب عموماً از داده‌های گوناگونی استفاده می‌کنند: در زبان‌شناسی صوری - زایشی از قضاوت‌های دستوری، همانند آزمون قضاوت دستوری که نشانگر دانش یا شم زبانی (Newmeyer, 2003: 682) است، استفاده می‌شود و در زبان‌شناسی نقش‌گرا عموماً داده‌های پیکره‌ای یا تجربی به‌کار می‌روند.

نیومیر (۲۰۰۳) بر این باور است که برای فهمیدن اینکه دستور یک زبان چطور کار می‌کند باید فراتر از داده‌های پیکره‌ای رفت؛ زیرا این داده‌ها بیانگر عملکرد زبانی هستند و نمی‌توانند تصویر روشنی از دستور یک زبان ارائه دهند. وی توصیه می‌کند که در پژوهش‌های نحوی به قضاوت‌های دستوری افراد که منعکس‌کننده شم زبانی آن‌هاست بیشتر توجه شود. او برای اثبات ادعای خود اظهار می‌دارد که برای پژوهشگرانی که در حوزه پژوهشی فراگیری زبان اول

فعالیت می‌کنند تفاوت میان شم زبانی و استفاده زبانی امری بارز و حل‌شده است. برای نمونه، هیرش پاشیک و گُلینگف^{۱۹} (1996) نشان دادند که بچه‌های ۱۳ و ۱۵ ماهه‌ای که در مرحله تک‌کلمه‌ای^{۲۰} هستند از این مسئله آگاهی دارند که کلمه‌هایی که به‌طور زنجیروار (یعنی بدون هیچ مکث و توقفی) به آن‌ها ارائه می‌شود شامل سازه‌های مجزا هستند. آن‌ها همچنین به این مسئله پی بردند که بچه‌های ۲۸ ماهه‌ای که تنها می‌توانند جمله‌های چهارکلمه‌ای تولید کنند قادرند با استفاده از موضوع فعل به معنی آن فعل پی ببرند. به‌یقین، پی بردن به این نکته‌ها درمورد درک زبانی بچه‌ها تنها با تحلیل تولیدات زبانی آن‌ها ممکن نبود. از این رو، نیومیر اظهار می‌دارد که دستور زبان پدیده‌ای کاملاً درونی است و برای بررسی ابعاد گوناگون آن نمی‌توان تنها منتظر نمود آن در تولیدات زبانی بود.

با وجود این، میر و تائو^{۲۱} (2005) در پاسخ به نیومیر (2003) این انگاشته را که داده‌های قضاوتی بیانگر شم زبانی هستند زیر سؤال بردند و استدلال کردند شم زبانی با دانش زبانی برابر نیست. آن‌ها در ادامه اظهار داشتند که به‌جای اینکه داده‌های قضاوتی را مهم‌تر تلقی کنیم باید آن‌ها را به‌صورت مکمل داده‌های پیکره‌ای و عملکردی ببینیم. به بیان دیگر، استفاده هم‌زمان و موازی از داده‌های قضاوتی و عملکردی و پیکره‌ای فهم ما را از پدیده مورد بررسی غنی‌تر می‌کند و به ما این امکان را می‌دهد تا تصویر کامل‌تری از پدیده مورد بررسی به‌دست آوریم. در همین راستا، پژوهش پیش‌رو تلاش می‌کند تا به‌وسیله یک آزمون قضاوت دستوری یافته‌های پژوهش‌های پیشین را تکمیل کند.

به‌طور خلاصه، تقریباً همه دانش ما از پیش‌سازندسازی گروه اسمی سنگین در زبان فارسی از پژوهش‌های پیکره‌ای به‌دست آمده است که منعکس‌کننده عملکرد زبانی است نه شم زبانی. آنچه باید به این ادبیات در حال رشد اضافه شود این است که شم گویشوران بومی زبان فارسی نسبت به پیش‌سازندسازی گروه اسمی سنگین چیست؟ ما در این پژوهش شواهدی ارائه می‌کنیم که نشان می‌دهند توالی مفعول مستقیم (چه کوتاه و چه بلند) و غیرمستقیم در جمله‌هایی که فارسی‌زبان‌ها تولید می‌کنند بسیار متفاوت از توالی‌ای است که آزمون قضاوت دستوری به ما می‌گوید؛ یافته‌ای که تنها با استفاده از داده‌های پیکره‌ای و عملکردی نمایان نمی‌شد.

۳. روش پژوهش

۳-۱. شرکت‌کننده‌ها

در آغاز، ۸۲ گویشور بومی زبان فارسی - که بین ۱۴ تا ۱۹ سال داشتند - در این پژوهش شرکت کردند. بعد از دادن اولین آزمون، بر اساس معیاری که در ادامه بیان خواهد شد تعدادی برای آزمون دوم گزینش شدند. آن‌ها دانش‌آموزان مقطع راهنمایی یا دبیرستان بودند و طبق پرسش‌نامه‌ای که در ابتدای پژوهش به این افراد داده شد یادگیری زبان انگلیسی آن‌ها تنها به کلاس‌های زبان مدرسه محدود می‌شد که نهایتاً هفته‌ای دو ساعت و از مقطع دوم راهنمایی شروع می‌شد. دانش‌آموزانی که بیش از ۳ سال در آموزشگاه‌های خصوصی انگلیسی آموخته بودند حذف شدند؛ زیرا در این صورت ممکن بود زبان فارسی آن‌ها تا حدی تحت تأثیر زبان انگلیسی قرار گیرد. علاوه بر این، هیچ کدام از این افراد قبل از شروع پژوهش در کشوری دیگر زندگی نکرده بودند و در خانه یا با دوستان به زبانی غیر از فارسی سخن نمی‌گفتند.

۳-۲. ابزار پژوهش

۳-۲-۱. آزمون یادآوری جمله

برای این آزمون در ابتدا ۲۰ دسته چهارجمله‌ای طراحی شد. در زیر، یکی از این دسته‌ها را می‌بینید:

جمله نوع اول: مفعول مستقیم کوتاه - مفعول غیرمستقیم (پرستار نمونه را به آزمایشگاه فرستاد)؛

جمله نوع دوم: مفعول غیرمستقیم - مفعول مستقیم کوتاه (پرستار به آزمایشگاه نمونه را فرستاد)؛

جمله نوع سوم: مفعول مستقیم بلند - مفعول غیرمستقیم (پرستار نمونه‌ای را که از بیمار گرفته بود به آزمایشگاه فرستاد)؛

جمله نوع چهارم: مفعول غیرمستقیم - مفعول مستقیم بلند (پرستار نمونه‌ای را که از بیمار گرفته بود به آزمایشگاه فرستاد).

در جمله‌های نوع اول و دوم - که فارسی‌زبانان به‌منظور بررسی توالی ترجیحی مفعول‌های مستقیم و غیرمستقیم طراحی کردند - از مفعول مستقیم و غیرمستقیم کوتاه، و در جمله‌های

نوع سوم و چهارم - که برای بررسی پیشایندسازی یا پسایندسازی گروه اسمی سنگین در فارسی طراحی شد - از مفعول مستقیم بلند و مفعول غیرمستقیم کوتاه استفاده شد. جمله‌های نوع اول و سوم دارای تناوب «مفعول مستقیم - مفعول غیرمستقیم» و جمله‌های نوع دوم و چهارم دارای تناوب «مفعول غیرمستقیم - مفعول مستقیم» بودند. برای تمرکز بیشتر بر هدف اصلی پژوهش و به دست آوردن نتایج بهتر همه مفعول‌های مستقیم از یک نوع یعنی از نوع مشخص با نشانگر «را» انتخاب شدند و تنها از فعل «فرستادن» که یک فعل دومفعولی با بسامد بالاست استفاده شد. همچنین، با اضافه کردن یک بند موصولی به مفعول مستقیم جمله‌های نوع سوم و چهارم دارای سازه سنگین شدند.

سپس با استفاده از این جمله‌ها، ۴۰ بلوک چهارجمله‌ای ساخته شد. هر بلوک شامل یک جمله با مفعول مستقیم کوتاه، یک جمله با مفعول مستقیم بلند، و دو جمله منحرف‌کننده^{۲۲} بود. در مرحله بعدی، هر ۱۰ بلوک یک نسخه از آزمون یادآوری جمله را تشکیل دادند، به گونه‌ای که هر شرکت‌کننده تنها یک جمله از چهار جمله تشکیل‌دهنده از هر دسته را می‌دید. دو بلوک تمرینی نیز به منظور آشنا کردن شرکت‌کنندگان با نوع آزمون طراحی شد و در نتیجه هر شرکت‌کننده یک نسخه ۱۲ بلوکی دریافت می‌کرد که ۲ بلوک اول آن تمرینی و ۱۰ بلوک دیگر بلوک‌های اصلی بودند. علاوه بر این، برای جلوگیری از تأثیر خستگی در بلوک‌های پایانی، توالی بلوک‌ها در هر نسخه برعکس گردید و به این ترتیب از هر نسخه دو حالت متفاوت به دست آمد. برای نمونه، در مورد نسخه ۱، دو حالت اول و دوم به صورت زیر طراحی شدند:

نسخه ۱ (حالت اول): بلوک تمرینی ۱ + بلوک تمرینی ۲ + بلوک ۱ + بلوک ۲ + ... + بلوک ۹ + بلوک ۱۰؛

نسخه ۱ (حالت دوم): بلوک تمرینی ۱ + بلوک تمرینی ۲ + بلوک ۱۰ + بلوک ۹ + ... + بلوک ۲ + بلوک ۱.

جمله‌ها با نرم‌افزار PowerPoint برای شرکت‌کنندگان به صورت انفرادی به صورت انفرادی ارائه می‌شد. زمان ارائه هر جمله نیز با توجه به طول جمله و همچنین نتایج اجرای مقدماتی که قبل از آغاز پژوهش از جمعیت مشابه صورت گرفته بود تعیین شد: جمله‌های با مفعول مستقیم بلند، جمله‌های با مفعول مستقیم کوتاه، و جمله‌های منحرف‌کننده هر کدام به ترتیب به مدت پنج، سه، و چهار ثانیه ارائه شدند. بعد از نمایش هر بلوک یک معادله ساده

ریاضی ارائه می‌شد که شرکت‌کنندگان می‌بایست درستی جواب آن را قضاوت می‌کردند. این کار مانع می‌شد که آن‌ها جمله‌ها را حفظ کنند و صرفاً به صورت طوطی‌وار بازگو کنند. بعد از قضاوت درستی معادله‌های ریاضی، آن‌ها می‌بایست با دیدن یک کلمه از هر جمله کل جمله را تکرار می‌کردند و این تکرار به دست پژوهشگر ضبط می‌شد.

در مرحله نمره دادن، شرکت‌کننده‌هایی که نتوانستند حداکثر ۸ معادله ریاضی را به درستی قضاوت کنند، کنار گذاشته شدند؛ زیرا این احتمال وجود داشت که آن‌ها بعد از دیدن جمله‌های هر بلوک به جای تمرکز بر قضاوت معادله‌های ریاضی جمله‌های دیده‌شده را در ذهن خود تکرار و حفظ کنند. علاوه بر این، آن‌هایی که نتوانستند از ۲۰ جمله منحرف‌کننده حداقل ۱۵ جمله را به یاد آورند از پژوهش کنار گذاشته شدند؛ زیرا این نشان می‌داد این افراد دقت و تمرکز کافی را در آزمون نداشتند. سپس برای نمره دادن به عملکرد شرکت‌کننده‌های باقی‌مانده، از مفهوم «معکوس‌سازی» استفاده شد؛ بدین معنی که تعداد جمله‌هایی که شرکت‌کننده‌ها با توالی مفعولی معکوس نسبت به آنچه روی صفحه نمایشگر دیده بودند تولید می‌کردند بر تعداد کل جمله‌های تولیدشده (جمله‌های عیناً تکرار شده + جمله‌های تولیدشده با توالی معکوس) تقسیم می‌شد و مقدار به دست آمده تبدیل به درصد می‌شد.

۲-۳. آزمون قضاوت دستوری

برای اینکه به ششم زبانی شرکت‌کنندگان درباره توالی مفعول‌های مستقیم و غیرمستقیم پی ببریم یک آزمون قضاوت دستوری کتبی بدون محدودیت زمانی طراحی کردیم. این آزمون از ۱۸ جفت جمله تشکیل شد: ۶ جفت جمله آزمایشی که در هر جفت توالی مفعول‌ها با هم تفاوت داشت (۳ جفت با مفعول مستقیم کوتاه، و ۳ جفت با مفعول مستقیم بلند)، و ۱۲ جفت جمله منحرف‌کننده. گنجاندن جمله‌های منحرف‌کننده در کنار جمله‌های آزمایشی به دو منظور انجام شد: اول اینکه شرکت‌کننده‌ها متوجه هدف اصلی آزمون نشوند، و دوم اینکه چون برای جمله‌های آزمایشی جواب درست یا غلط وجود نداشت، در نظر گرفتن عملکرد شرکت‌کننده‌ها در قضاوت جمله‌های منحرف‌کننده می‌توانست نشان دهد که آن‌ها تا چه اندازه با تمرکز و دقت آزمون را انجام دادند. به همین دلیل ابتدا نمره افراد در جمله‌های منحرف‌کننده محاسبه می‌شد و اگر آن‌ها ۸۰ درصد این جمله‌ها را درست قضاوت می‌کردند، نمره آن‌ها برای جمله‌های

آزمایشی محاسبه می‌شد. درمورد جمله‌های آزمایشی نیز شرکت‌کنندگان باید این جمله‌ها را می‌خواندند و به آن‌ها از ۱ (خیلی بد) تا ۶ (خیلی خوب) نمره می‌دادند.

۳-۳. روش انجام پژوهش

تمام مراحل گردآوری داده‌ها در یکی از شعبه‌های کانون زبان ایران انجام گرفت. در آغاز که با زبان‌آموزانی که زودتر از زمان کلاس خود در مؤسسه حاضر می‌شدند و همچنین در سطح مبتدی یادگیری زبان انگلیسی بودند صحبت می‌شد و چنانچه ایشان به شرکت در پژوهش ابزار تمایل می‌کردند پرسش‌نامه کوتاهی درمورد میزان و نحوه آشنایی آن‌ها با زبان‌هایی غیر از زبان انگلیسی داده می‌شد. بعد از پر کردن این پرسش‌نامه به آن‌ها آزمون اول یعنی آزمون عملکردمحور یادآوری جمله داده شد. سپس از میان ۸۲ فارسی‌زبان که این آزمون را انجام دادند ۳۸ نفر (طبق ملاک‌هایی که در بالا گفته شد) برای آزمون بعدی یعنی آزمون قضاوت دستوری دعوت شدند.

آزمون قضاوت دستوری حداقل یک هفته بعد از آزمون یادآوری جمله به شرکت‌کننده‌ها داده شد. این آزمون به‌صورت انفرادی یا گروه‌های چندنفره اجرا شد و در ابتدا به شرکت‌کنندگان به زبان فارسی توضیحاتی درمورد چگونگی انجام آن ارائه شد. در مرحله نمردهی آزمون قضاوت دستوری برای اطمینان از اینکه شرکت‌کنندگان با دقت این آزمون را انجام داده‌اند، تنها شرکت‌کنندگانی که توانسته بودند حداقل ۱۹ جمله (۸۰ درصد) از ۲۴ جمله منحرف‌کننده (که دارای پاسخ درست و غلط بودند) را به‌درستی قضاوت و نمره‌گذاری کنند در پژوهش باقی ماندند. از این رو، یک شرکت‌کننده که تنها ۱۸ جمله منحرف‌کننده را به‌درستی قضاوت کرده بود کنار گذاشته شد و در نهایت داده‌های مربوط به ۳۷ شرکت‌کننده به مرحله تحلیل آماری رسید.

۴. یافته‌ها

۴-۱. آزمون یادآوری جمله

هدف از آزمون این بود که پی ببریم فارسی‌زبانان - وقتی از آن‌ها خواسته می‌شود جمله‌هایی را که از قبل خوانده‌اند تکرار کنند - مفعول مستقیم کوتاه و بلند را در کجای جمله قرار می‌دهند.

برای یافتن پاسخ این سؤال بعد از اینکه میزان معکوس‌سازی تناوب مفعول‌ها تبدیل به درصد شدند از آزمون آنالیز واریانس با اندازه‌گیری‌های مکرر با طراحی ۲×۲ (دو جایگاه متفاوت برای مفعول مستقیم × دو طول متفاوت برای مفعول مستقیم) استفاده شد.

همان‌طور که جدول ۱ آمار توصیفی مربوط به این آزمون را نشان می‌دهد تمام جمله‌های نوع اول و سوم، یعنی جمله‌هایی که دارای تناوب «مفعول مستقیم - مفعول غیرمستقیم» بودند (چه مفعول مستقیم کوتاه و چه مفعول مستقیم بلند)، با همین تناوب از طریق شرکت‌کنندگان یادآوری شدند؛ اما بیشتر جمله‌های نوع دوم و چهارم، یعنی جمله‌هایی که با تناوب «مفعول غیرمستقیم - مفعول مستقیم» ارائه شدند (چه مفعول مستقیم کوتاه و چه مفعول مستقیم بلند)، با تناوب معکوس یعنی «مفعول مستقیم - مفعول غیرمستقیم» به یاد آورده شدند (۹۳ درصد جمله‌های دارای مفعول مستقیم کوتاه و ۸۷ درصد جمله‌های دارای مفعول مستقیم بلند).

جدول ۱: میانگین درصد جمله‌هایی که با تناوب معکوس به یاد آورده شدند

Table 1: Descriptive statistics for the proportion of inverted sentences in each sentence type

ترتیب مفعولی		مستقیم - غیرمستقیم		غیرمستقیم - مستقیم	
طول مفعول مستقیم		کوتاه	بلند	کوتاه	بلند
نوع جمله		اول	سوم	دوم	چهارم
میانگین (%)		۰	۰	۹۳	۸۷
انحراف معیار		۰	۰	۰/۲۳	۰/۲۸

نتایج آزمون آنالیز واریانس با اندازه‌گیری‌های مکرر نشان داد که متغیر جایگاه مفعول مستقیم تأثیر معناداری بر میزان معکوس‌سازی تناوب مفعول‌ها داشته است:

$$F(1, 36) = 669.32, p = .000$$

متغیر طول مفعول مستقیم تأثیر معناداری بر میزان معکوس‌سازی تناوب مفعول‌ها نداشته است:

$$F(1, 36) = 1.37, p = .24$$

در ضمن هیچ تقابلی بین متغیرهای طول مفعول و جایگاه مفعول دیده نشد:

$$F(1, 36) = .86, p = .36$$

۴-۲. آزمون قضاوت دستوری

در این آزمون، شرکت‌کنندگان می‌بایست از ۱ تا ۶ به جمله‌های نوع اول تا چهارم نمره می‌دادند. جدول ۲ بسامد هر یک از نمره‌هایی را که به این جمله‌ها داده شد نشان می‌دهد. این نکته را نیز باید یادآور شد که هر کدام از ۳۷ شرکت‌کننده این پژوهش ۳ نمونه (۳×۳۷) از هر ۴ نوع جمله (با طول و تناوب متفاوت مفعول‌ها) را ارزیابی کردند.

جدول ۲: بسامد نمره‌های داده‌شده به هر کدام از چهار نوع جمله

Table2: The frequency of each score given to each sentence type in the GJT

ترتیب مفعولی		مستقیم - غیر مستقیم		غیر مستقیم - مستقیم	
طول مفعول مستقیم		بلند	کوتاه	بلند	کوتاه
نوع جمله		اول	سوم	دوم	چهارم
۱	خیلی بد	۰	۰	۸	۱۴
۲	بد	۰	۰	۶	۷
۳	نسبتاً بد	۰	۰	۱۳	۱۷
۴	نسبتاً خوب	۰	۲	۱۶	۱۱
۵	خوب	۱	۱۰	۳۴	۳۰
۶	خیلی خوب	۱۱۰	۹۹	۳۴	۳۲
مجموع		۱۱۱	۱۱۱	۱۱۱	۱۱۱

۴-۲-۱. بررسی جمله‌های دارای مفعول مستقیم کوتاه

همان‌طور که جدول ۲ نشان می‌دهد تقریباً همه جمله‌های نوع اول نمره ۶ گرفتند (به‌استثنای یک نمره ۵). ولی نمره‌هایی که به جمله‌های نوع دوم داده شد پراکندگی نسبتاً زیادی داشت: نمره‌های ۵ و ۶ هر کدام ۳۴ بار به این جمله‌ها داده شدند؛ در حالی که نمره‌های ۱، ۲، ۳، و ۴ هر کدام به ترتیب ۸، ۶، ۱۳، و ۱۶ بار برای این جمله‌ها انتخاب شدند.

برای اینکه از اهمیت آماری اختلاف بین این بسامدها آگاه شویم بر آن شدیم تا نمره‌ها را به دو دسته قابل قبول (نمره‌های ۴، ۵، و ۶) و غیر قابل قبول (نمره‌های ۱، ۲، و ۳) تقسیم کنیم. به بیان دیگر، اگر شرکت‌کننده‌ها برای جمله‌ای عدد ۴ و بالاتر از آن را انتخاب می‌کردند، به این

معنی بود که آن جمله از نظر آن‌ها قابل قبول است و اگر برای جمله‌ای اعداد ۳ و پایین‌تر از آن را انتخاب می‌کردند، به این معنی بود که آن‌ها آن جمله را غیر قابل قبول می‌پنداشتند. جدول ۳ نشان می‌دهد چه تعداد از جمله‌های نوع اول تا چهارم از شرکت‌کنندگان نمره‌های قابل قبول و غیر قابل قبول دریافت کردند.

جدول ۳: تعداد جمله‌هایی که از نظر شرکت‌کنندگان قابل قبول و غیر قابل قبول بود
Table3: Frequency of the sentences rated as "unacceptable" and "acceptable"

ترتیب مفعولی	مستقیم - غیر مستقیم	غیر مستقیم - مستقیم		
طول مفعول مستقیم	بلند	کوتاه	بلند	کوتاه
نوع جمله	اول	سوم	دوم	چهارم
قابل قبول	۱۱۱	۱۱۱	۸۴	۷۳
غیر قابل قبول	۰	۰	۲۷	۳۸

درمورد جمله‌های نوع اول نیازی به آزمون نبود؛ زیرا همه آن‌ها در دسته قابل قبول بودند. ولی درمورد جمله‌های نوع دوم نتیجه آزمون کای مجذور نشان داد که بین تعداد جمله‌هایی که نمره قابل قبول گرفتند و تعداد جمله‌هایی که نمره غیر قابل قبول گرفتند تفاوت معنی‌داری وجود دارد: $\chi^2(1, 111) = 29.27, p = .000$

۲-۲-۴. بررسی جمله‌های دارای مفعول مستقیم بلند

همان‌طور که در جدول ۲ می‌توان دید جمله‌های نوع سوم ۹۹ بار نمره ۶، ۱۰ بار نمره ۵، و ۲ بار نمره ۴ گرفتند و هیچ کدام از شرکت‌کنندگان به آن‌ها نمره‌های ۱، ۲، یا ۳ ندادند؛ اما درمورد جمله‌های نوع چهارم نمره‌های داده‌شده پراکندگی بیشتری دارند. این جمله‌ها ۳۲ بار نمره ۶، ۳۰ بار نمره ۵، ۱۱ بار نمره ۴، ۱۷ بار نمره ۳، ۷ بار نمره ۲، و ۱۴ بار نمره ۱ دریافت کردند.

همانند جمله‌های نوع اول و دوم نمره‌های داده‌شده به این جمله‌ها نیز به دو دسته قابل قبول (نمره‌های ۴، ۵، و ۶) و غیر قابل قبول (نمره‌های ۱، ۲، و ۳) تقسیم شدند. همان‌طور که در جدول ۳ می‌توان دید همه جمله‌های نوع سوم نمره قابل قبول گرفتند و بنابراین نیازی به آزمون کای مجذور نبود. ولی جمله‌های نوع چهارم ۷۳ بار نمره قابل قبول و ۳۸ بار نمره غیر قابل قبول

گرفتند و نتیجه آزمون کای مجذور نشان داد که بین این جمله‌ها تفاوت معنی‌داری وجود دارد:
 $\chi^2(1, 111) = 12.33, p = .000$
 به‌طور خلاصه تمام جمله‌هایی که تناوب «مفعول مستقیم - مفعول غیرمستقیم» داشتند، چه مفعول مستقیم بلند و چه مفعول مستقیم کوتاه، از شرکت‌کنندگان نمره قابل قبول گرفتند؛ ولی درمورد جمله‌های با تناوب «مفعول غیرمستقیم - مفعول مستقیم» (چه کوتاه و چه بلند) نمره‌های شرکت‌کنندگان پراکنده بود که در نهایت نتایج آزمون‌های کای مجذور نشان داد که این جمله‌ها نیز برای شرکت‌کنندگان به میزان معناداری قابل قبول‌اند.

۵. بحث و تفسیر یافته‌ها

پژوهش حاضر با هدف بررسی جایگاه مفعول مستقیم در زبان فارسی و چگونگی پیش‌اندسازی یا پس‌اندسازی مفعول مستقیم سنگین در این زبان انجام شد. برای رسیدن به این هدف از دو آزمون متفاوت استفاده شد: یکی آزمون یادآوری جمله که برای گردآوری داده‌های عملکردمحور و برای مقایسه با داده‌های مشابه پژوهش‌های پیشین طراحی شد، و دیگری آزمون قضاوت دستوری که برای پی بردن به شم زبانی (Newmeyer, 2003) گویشوران بومی زبان فارسی درمورد پدیده مورد بررسی طراحی شد.

در آزمون یادآوری جمله - که در آن شرکت‌کنندگان می‌بایست جمله‌ها را بر روی نمایشگر رایانه می‌دیدند و بعد از قضاوت درستی یک معادله ریاضی ساده آن‌ها را با کمک یک واژه از همان جملات که به ایشان ارائه شده بود به‌یاد می‌آوردند - آن‌ها تناوب مفعولی هیچ کدام از جمله‌های نوع اول و سوم را تغییر ندادند. در عوض، بیشتر جمله‌های نوع دوم و چهارم (به‌ترتیب ۹۳ و ۸۷ درصد) را با تناوب مفعولی معکوس یعنی «مفعول مستقیم - مفعول غیرمستقیم» به‌خاطر آوردند و بنابراین فرضیه نخست این پژوهش تأیید شد. نکته مهم دیگر این است که قضاوت درستی معادله‌های ریاضی بعد از دیدن جمله‌ها به شرکت‌کنندگان اجازه نمی‌داد که فقط از حفظ و طوطی‌وار تکرار کنند؛ بلکه آن‌ها مجبور بودند با استفاده از مفهوم جمله که در ذهنشان مانده بود جمله‌ها را از نو بسازند (Bock & Warren, 1985). این موضوع این امکان را به ما داد تا تناوب کلمه‌ای ترجیحی شرکت‌کنندگان را در هنگام تولید زبان بیابیم.

ولی یافته جالب توجه پژوهش پیش رو این بود که در آزمون قضاوت دستوری، خلاف آنچه در آزمون یادآوری جمله و همچنین پژوهش‌های پیشین (Faghiri & Samvelian, 2014; Faghiri & et al., 2014) دیده شد، شرکت‌کنندگان هر دو تناوب «مفعول مستقیم - مفعول غیرمستقیم» و «مفعول غیرمستقیم - مفعول مستقیم» را قابل قبول دانستند (همه ۱۱۱ جمله نوع اول و ۸۴ جمله از ۱۱۱ جمله نوع دوم نمره قابل قبول دریافت کردند). حتی در مورد جمله‌های دارای مفعول مستقیم بلند باز هم هر دو تناوب «مفعول مستقیم - مفعول غیرمستقیم» و «مفعول غیرمستقیم - مفعول مستقیم» نمره قابل قبول گرفتند (همه ۱۱۱ جمله نوع سوم و ۷۳ عدد از ۱۱۱ جمله نوع چهارم). از این نتایج چنین برمی‌آید که فرضیه دوم پژوهش در مورد قابل قبولی هر دو توالی «مفعول مستقیم - مفعول غیرمستقیم» و «مفعول غیرمستقیم - مفعول مستقیم» برای گویشوران بومی فارسی تأیید می‌شود. به نظر می‌رسد که با توجه به تناوب کلمه‌ای نسبتاً آزاد زبان فارسی ششم زبانی فارسی‌زبانان هر دو تناوب بررسی شده را قابل قبول می‌داند. البته در این پژوهش جمله‌ها به صورت مجزا و به دور از بافت ارائه شدند؛ شاید اگر جمله‌ها در بافت معینی قرار گیرند، عوامل دیگری مانند تأکید^{۲۳} و مبتدا^{۲۴} نیز در تناوب مفعول‌ها اثرگذار باشد (Karimi, 2003; Moinzadeh, 2001). ولی از آنجا که در این پژوهش از شرکت‌کنندگان خواسته شد صرفاً از لحاظ دستوری و تنها با تکیه بر ششم زبانی خود جمله‌ها را مورد قضاوت قرار دهند، آن‌ها هر چهار نوع جمله را قابل قبول دانستند.

نتایج به دست آمده از آزمون یادآوری جمله با یافته‌های پیکره‌ای کریمی (2003) و یافته‌های تجربی فقیری و سامولیان (2014) و فقیری و همکاران (2014) همخوانی دارد. پژوهشگران نام‌برده اظهار کرده بودند که در فارسی مفعول مشخص با نشانگر «را» قبل از مفعول غیرمستقیم و دور از فعل قرار می‌گیرد؛ اما نکته‌ای که اهمیتی دوچندان دارد این است که در آزمون یادآوری جمله شرکت‌کنندگان بیشتر جمله‌های نوع دوم و چهارم را به جمله‌های نوع اول و سوم تغییر دادند (یعنی جمله‌ها را با تناوب «مفعول مستقیم - مفعول غیرمستقیم» به خاطر آوردند)؛ ولی در آزمون قضاوت دستوری آن‌ها جمله‌های نوع دوم و چهارم را مانند جمله‌های نوع اول و سوم (هرچند به میزان نسبتاً پایین‌تر) قابل قبول دانستند. این یافته‌ها نشان می‌دهند که در پژوهش‌های نحوی نباید فقط به داده‌های پیکره‌ای و عملکردی بسنده کرد؛ بلکه باید از داده‌های درون‌نگرانه^{۲۵} و قضاوتی نیز بهره جست (Newmeyer, 2003). در همین راستا باد،

هی، و جابندی^{۲۶} (2003) اظهار می‌کنند که زبان یک سیستم احتمالی^{۲۷} است که در آن قوانین دستوری بیشتر از اینکه دوقطبی باشند دارای درجه‌بندی هستند (ر. ک: پی‌نوشت‌های ۱۶ و ۱۸). در نتیجه دانش زبانی را نباید به‌صورت مجموعه‌ای از یک سری قوانین یا محدودیت‌های دوقطبی دانست؛ بلکه باید آن را به‌صورت مجموعه‌ای از قوانین درجه‌بندی‌شده در نظر گرفت. به همین دلیل، تصویری که داده‌های تجربی، عملکردمحور، و پیکره‌ای از دانش زبانی افراد به ما می‌دهند تصویر کاملی نیست و می‌بایست با داده‌های درون‌نگرانه و قضاوتی تکمیل شود تا ماهیت درجه‌بندی‌شده دانش زبانی افراد به شکل بهتری نمایان شود.

در مجموع، از نظر پیشایندسازی گروه اسمی سنگین در فارسی تصویری که از کنار هم قرار دادن داده‌های حاصل از آزمون یادآوری جمله و آزمون قضاوت دستوری به‌دست می‌آوریم زبانی را به ما نشان می‌دهد که همانند زبان ژاپنی بیشتر مفهوم‌محور است تا صورت‌محور (Chang, 2009; Yamashita & Chang, 2001). در عوض، زبانی مانند انگلیسی که تناوب کلمه‌ای آزادی ندارد و در آن موضوع‌های فعل بعد از فعل می‌آیند، بیشتر به بُعد صوری تمایل دارد. علاوه بر این، باید یادآور شد که اصل شناسایی سریع سازه‌های بلافصل هاوکینز (1994) - که به‌روشنی پسایندسازی گروه اسمی سنگین را در انگلیسی و پیشایندسازی آن را در ژاپنی توجیه می‌کند - نمی‌تواند این پدیده را در فارسی توضیح دهد. از این رو، بهترین توضیح برای جایگاه مفعول مستقیم و پیشایندسازی گروه اسمی سنگین در فارسی این است که در این زبان تناوب مفعول‌ها در محدوده قبل از فعل به میزان مشخص یا معین بودن مفعول مستقیم بستگی دارد.

۶. نتیجه‌گیری

در این پژوهش بر آن شدیم تا با گردآوری داده‌های قضاوت‌محور در کنار داده‌های عملکردمحور از گویشوران بومی زبان فارسی یافته‌های پژوهش‌های پیشین درمورد تناوب مفعول‌ها و پیشایندی گروه اسمی سنگین در فارسی را کامل کنیم. انگیزه ما از انجام چنین مطالعه‌ای کامل کردن تصویری بود که پژوهش‌های پیشین (که غالباً از داده‌های عملکردمحور و پیکره‌محور استفاده نموده بودند) از پدیده مورد بررسی داده بودند (Wasaw & Arnold,

2003). این تصویر کامل‌تر زبانی را نشان می‌دهد که اثرپذیری بسیار بیشتری نسبت به عوامل مفهومی دارد؛ زبانی که در آن میزان مشخص بودن مفعول مستقیم نقش مهمی در جایگاه آن در جمله دارد؛ زبانی مشابه زبان‌هایی با تناوب کلمه‌ای آزاد مانند ژاپنی و کره‌ای از نظر تأثیرپذیری از بُعد مفهومی ولی متفاوت از آن‌ها از این جهت که تمایل «بلند قبل از کوتاه» یا پیشایندسازی گروه اسمی سنگین در آن همیشه دیده نمی‌شود.

در پایان، لازم می‌دانیم پیشنهادهایی برای پژوهش‌های آتی ارائه دهیم. اولین پیشنهاد این است که در این پژوهش‌ها در آزمون‌های قضاوت دستوری عوامل معنایی و منظورشناسی را در جمله‌ها لحاظ کنند (به بیان دیگر جمله‌ها در بافت ارائه شوند) و در جمله‌ها میزان مشخص بودن مفعول مستقیم را کنترل کنند. همچنین پیشنهاد می‌کنیم در این آزمون‌ها از فعل‌های دومفعولی دیگر (Pinker, 1989) و گروه‌های حرف اضافه‌ای که دارای حرف اضافه‌های دیگری هستند استفاده شود. علاوه بر آزمون‌های قضاوت دستوری گزارش‌های زبانی درون‌نگرانه و عقب‌نگرانه^{۲۸} (Rebuschat, 2013) مانند تکنیک «برون‌فکنی اندیشه»^{۲۹} می‌توانند درچه‌های دیگری به سوی بررسی شم زبانی گویشوران زبان‌های بومی بگشایند. پیشنهاد پایانی نیز مربوط می‌شود به نکته‌ای درباره تولیدات زبانی و داده‌های پیکره‌ای که در این پژوهش به صورت نقطه مقابل شم زبانی از آن صحبت شد: حدود نیم قرن پیش، راس (1967: 49) ادعا کرد که «تمام بحث ما درباره اینکه چه گروه‌های اسمی‌ای سنگین یا پیچیده هستند به شیوه استفاده از زبان منتهی می‌شود، و اینکه گستره وسیعی از متغیرهای گویشی و حتی گویش فردی^{۳۰} در سنگینی یا پیچیدگی این سازه‌ها نقش دارند». بنابراین، اکیداً توصیه می‌شود که برای به دست آوردن تصویری کامل‌تر از پدیده جابه‌جایی گروه اسمی سنگین در فارسی پژوهش‌های آتی از پیکره‌هایی استفاده کنند که بهتر نمایانگر زبان فارسی طبیعی و محاوره‌ای باشند.

۷. پی‌نوشت‌ها

۱. منظور از توالی بی‌نشان آن توالی‌ای است که بسامد بیشتری دارد، کاربرد اطلاعاتی خاصی را منظور نمی‌کند، و درواقع می‌تواند برای هر ساختار اطلاعاتی به‌کار رود (Faghiri & Samvelian, 2014).
2. The end-weight principle
3. Incremental models of sentence production
۴. این پدیده در انگلیسی با نام Heavy NP Shift شناخته می‌شود و ما در این پژوهش این پدیده را بر حسب جهت حرکت سازه سنگین «پیشایندسازی» یا «پسایندسازی» سازه سنگین می‌نامیم.
۵. معین‌زاده (2001) و مرعشی (1970) با ارائه مدرک‌هایی استدلال می‌کنند که همه گروه‌ها (حتی گروه فعلی) در فارسی هسته - آغازند، و تناوب کلمه‌ای این زبان، همانند انگلیسی، فاعل - فعل - مفعول است.
6. scrambling
7. Faghiri and Samvelian
8. judgement data
9. Heaviness
10. Chomsky
11. embedded clause
12. accessibility
13. grammatical encoding
14. differential object marking (DOM) criterion
15. formal-generative linguistics
۱۶. دوقطبی بودن یا categoricity در این مکتب به این باور اشاره دارد که جمله‌ها در بررسی‌های نحوی یا درست‌اند یا نادرست.
17. functional linguistics
۱۸. منظور از درجه‌بندی یا gradience این است که در بررسی‌های نحوی جمله‌ها می‌توانند درجه‌های گوناگونی از درست بودن داشته باشند. به بیان دیگر، جمله‌ای که دستور زبان آن را کاملاً نادرست می‌پندارد می‌تواند در شرایطی خاص به دست بعضی افراد به دلایل گوناگون به‌کار برود و در پیکره‌های زبانی ثبت شود.
19. Hirsh-Pasek and Golinkoff
20. one-word stage
21. Meyer and Tao
22. distractor sentences
23. focus
24. topic
25. introspection data

26. Bod, Hay and Jannedy
27. probabilistic system
28. introspective and retrospective verbal reports
29. think-aloud protocols
30. idiolect

۸. منابع

- راسخ مهند، محمد (۱۳۸۳). «جایگاه مفعول مستقیم در فارسی». *نامه فرهنگستان*. د ۶. ش ۴ (پیاپی ۲۴). صص ۵۶-۶۶.
- راسخ مهند، محمد و مریم قیاسوند (۱۳۹۲). «عوامل مؤثر بر پسایندسازی در زبان فارسی». *زبان‌شناسی و گویش‌های خراسان*. د ۵. ش ۲ (پیاپی ۹). صص ۲۷-۴۷.
- _____ (۱۳۹۳). «بررسی پیکره‌بنیاد تأثیر عوامل نقشی در قلب نحوی کوتاه فارسی». *دستور*. ش ۱۰. صص ۱۶۳-۱۹۷.
- راسخ مهند، محمد و دیگران (۱۳۹۱). «تبیین نقشی خروج بند موصولی در زبان فارسی». *پژوهش‌های زبان‌شناسی*. د ۴. ش ۱. صص ۲۱-۴۰.
- شیخ‌الاسلامی، افتخار (۱۳۸۷). «نقش ساختار اطلاعاتی در خروج‌بندهای موصولی در زبان فارسی». پایان‌نامه کارشناسی ارشد. دانشگاه کردستان.
- علائی، مجید و دیگران. (۱۳۹۶). «چیدمان سازه‌ها در زبان فارسی متأثر از وزن دستوری: تبیینی پردازش‌محور». *جستارهای زبانی*. صص ۱-۲۹.
- مویدی، مونا و احمدرضا لطفی (۱۳۹۲). «بررسی ساخت دومفعولی در متون ادب فارسی». *پژوهش‌های زبان‌شناسی*. د ۵. ش ۱. صص ۱۰۱-۱۱۹.
- میردهقان، مهین‌ناز و سعیدرضا یوسفی (۱۳۹۵). «حرف اضافه‌نمایی افتراقی در وفسی در چهارچوب نظریه بهینگی». *جستارهای زبانی*. د ۷. ش ۳ (پیاپی ۳۱). صص ۱۹۷-۲۲۲.

References:

- Alaei, M.; M. Tehrani Doost & M. Rasekh-Mahand, (2017), "Constituent ordering in Persian under the influence of grammatical weight: A processing-based explanation". *Language Related Research*. Pp. 1-29. [In Persian].

- Arnold, J. E.; T. Wasow; A. Losongco & R. Ginstrom (2000). "Heaviness vs. newness: The effects of complexity and information structure on constituent ordering". *Language*. 76(1). Pp. 28-55.
- Bock, K. & R. Warren (1985). "Conceptual accessibility and syntactic structure in sentence formulation". *Cognition*. 21. Pp. 47-67.
- ----- & W. Levelt (1994). "Language Production: Grammatical Encoding". In M. A. Gernsbacher (Ed.), *Handbook of Psycholinguistics* (pp. 945-984). New York: Academic Press.
- ----- (1982). "Towards a cognitive psychology of syntax: information processing contributions to sentence formulation". *Psychological Review*. 89.Pp. 1-47.
- Bod, R., Hay, J. & S. Jannedy (2003). "Introduction". In R. Bod, J. Hay, and S. Jannedy (Eds.) *Probabilistic Linguistics* (pp. 1-10). Cambridge, MA: MIT Press.
- Chang, F. (2009). "Learning to order words: A connectionist model of heavy NP shift and accessibility effects in Japanese and English". *Journal of Memory and Language*, 61. Pp. 374-97.
- Choi, Hye-Won (2007). "Length and order: A corpus study of Korean dative-accusative construction". *Discourse and Cognition*, 14(3).Pp. 207-227.
- Chomsky, N. (1975). *The Logical Structure of Linguistic Theory*. New York: Plenum press.
- Faghiri, P. & P. Samvelian (2014). "Constituent ordering in Persian and the weight factor". In C. Pinon (Ed.), *Empirical issues in syntax and semantics 10* (EISS10), (pp. 215-232).
- ----- B. Hemforth (2014). "Accessibility and Word Order: The Case of Ditransitive Constructions in Persian". In S. Müller (éd.), *Proceedings of the 21st International Conference on Head-Driven Phrase Structure Grammar*. Stanford: CSLI Publications, pp. 217-237
- Frazier, L. & J. D. Fodor (1978). "The sausage machine: A two-stage parsing model". *Cognition*. 6. Pp. 291-325.

- Garrett, F. (1980). "Levels of processing in sentence production". *Language Production I*. Pp. 177-220.
- Hawkins, J. (1994). *A Performance Theory of order and Constituency*. Cambridge: Cambridge University Press.
- Hirsh-Pasek, K. & R. Golinkoff, (1996), *The origins of Grammar: Evidence from Early Language Comprehension*. Cambridge, MA: MIT Press.
- Karimi, S. (2003). "On scrambling in Persian". In S. Karimi (Ed.), *Word order and Scrambling* (pp. 301-324). Malden, MA: Blackwell Publishing.
- Kimball, J. (1973). "Seven principles of surface structure parsing in natural language". *Cognition* 2. Pp. 15-47.
- Levelt, W. J. M. (1989). *Speaking: From Intention to Articulation*. Cambridge: MIT Press.
- Lotfi, A. & M. Moayedi, (2013), "Double object construction in the Persian literary texts". *Researches in Linguistics*. 5(8). Pp. 101-119. [In Persian].
- Manning, C. (2003). "Probabilistic Syntax". In R. Bod, J. Hay, and S. Jannedy (Eds.), *Probabilistic Linguistics* (pp. 289-341). Cambridge, MA: MIT Press.
- Marashi, M. (1970). *The Persian Verb: A Partial Description for Pedagogical Purposes*. Unpublished Phd Dissertation, The University of Texas- Austin.
- Mazurkewich, I. (1984). "The acquisition of the dative alternation by second language learners and linguistic theory". *Language Learning* 34(1). Pp. 91-109.
- Meyer, C. & H. Tao, (2005), "Response to Newmeyer's 'Grammar is grammar and usage is usage'". *Language*. 81(1). Pp. 226-228.
- Mirdehghan, M. & S. Yusofi, (2016). "Differential adpositional case marking in Vafsi within Optimality Theory". *Language Related Research*. 7(3). Pp. 197-222. [In Persian].
- Moinszadeh, A. (2001). *An Antisymmetric, Minimalist Approach to Persian Phrase Structure*. Unpublished doctoral dissertation, University of Ottawa-Canada.

- Newmeyer, F. (2003). "Grammar is grammar and usage is usage". *Language*. 79.Pp.682-707.
- Pinker, S. (1989). *Learnability and Cognition: The Acquisition of Argument Structure*. Cambridge, MA: MIT Press.
- Rasekh-Mahand, M. & M. Ghiasvand, (2013), "Motivating factors of postposing in Persian". *Journal of Linguistics & Khorasan Dialects*, 5(9). Pp. 27-47. [In Persian].
- -----, (2014), "A Corpus-based study of functional motivations effects on Persian scrambling". *Grammar*. 10.Pp. 163-197. [In Persian].
- ----- (2004). "The position of direct object in Persian". *The Journal of the Persian Academy*, 6(4), pp. 55-66. [In Persian].
- Rasekh-Mahand, M.; M. Alizadeh Sahraie; R. Izadifar & M. Ghiasvand, (2012). "The functional explanation of relative clause extraposition in Persian". *Researches in Linguistics*. 4(1). Pp. 21-40. [In Persian].
- Rebuschat, P. (2013). "Measuring implicit and explicit knowledge in second language research". *Language Learning*. 63(3).Pp. 595-626.
- Ross, J. R. (1967). *Constraints on Variables in Syntax*. Unpublished Phd Dissertation, Massachusetts Institute of Technology- Massachusetts.
- Sheikholeslami, E. (2008). *The Role of Information Structure in Relative Clause Extraposition in Persian* (unpublished M.A Thesis). University of Kurdistan-Iran. [In Persian].
- Stallings, L. M.; M. C. MacDonald & P. G. O'seaghdha, (1998), "Phrasal ordering constraints in sentence production: Phrase length and verb disposition in heavy-NP shift". *Journal of Memory and Language*. 39(3). Pp.392-417.
- Wasaw, T. & J. Arnold, (2003), "Post-verbal constituent ordering in English". In B. Kortmann, and E. C. Traugott (Eds.), *Determinants of grammatical variation in English* (pp. 119-154). Berlin: Mouton de Gruyter.

- ----- (1997). "Remarks on grammatical weight". *Language Variation and Change*. 9. Pp.81-105.
- Yamashita, H. & F. Chang, (2001), "Long before short" preference in the production of a head-final language". *Cognition*. 81(2). Pp.45-B55.