

بهبود و توسعه یک سامانه مترجم‌یار انگلیسی به فارسی

زینب وکیل و شهرام خدیوی

آزمایشگاه فناوری زبان طبیعی، دانشکده مهندسی کامپیوتر و فناوری ارتباطات،

دانشگاه صنعتی امیرکبیر، تهران، ایران

چکیده

در این مقاله، برای اولین بار به بررسی سامانه‌های مترجم‌یار تعاملی و ارائه روش‌هایی در جهت بهبود کارایی این سامانه‌ها در زمینه ترجمه انگلیسی به فارسی می‌پردازیم. در یک سامانه مترجم‌یار تعاملی، ترجمه یک متن از یک سری همکاری‌های متناوب ماشین و مترجم انسانی به وجود می‌آید. مشارکت ماشین تنها در قالب پیشنهادهایی درباره جمله مقصد انجام می‌شود؛ اما مترجم می‌تواند آزادانه پیشنهادهای ارائه شده توسط مشارکت سامانه را بپذیرد، تغییر دهد و یا از آنها صرفه نظر کند. در واقع یک سامانه مترجم‌یار به عنوان ابزاری جهت تسهیل و تسریع فرآیند ترجمه به خدمت مترجم در می‌آید. ما در این مقاله با اشاره به نواقص روش‌های قبلی در بخش تعاملی سامانه مترجم‌یار، روش‌های جدیدی را ارائه داده‌ایم که در نهایت موجب بهبود ۱/۳ درصدی نتایج سامانه می‌شود.

واژگان کلیدی: سامانه مترجم‌یار رایانه‌ای (Computer-Assisted (Aided Translation)، سامانه مترجم‌یار تعاملی (Interactive CAT)، جستجوی پیشوند (Prefix Search).

۱- مقدمه

نیاز ترجمه متون از یک زبان به زبان دیگر، همواره یک نیاز اساسی و مهم بوده است. این نیاز، در بخش‌هایی نظیر روابط بین‌الملل و خبرگزاری‌ها به یک نیاز دائم و ضروری و در عین حال پرهزینه تبدیل می‌شود. برای مثال هزینه ترجمه برای تمام نهادهای اتحادیه اروپا در سال ۲۰۰۶، حدود هشتصد میلیون یورو برآورد شد، که از این مبلغ سیصد میلیون یورو مربوط به ترجمه متون بوده است. هزینه‌های بالای ترجمه انسانی متون، نهادهای بین‌المللی بزرگ نظیر اتحادیه اروپا را بر آن داشت تا روی پروژه‌های ترجمه خودکار متون، سرمایه‌گذاری نمایند؛ تحقیقات و مطالعات انجام شده در این پروژه‌ها که تاکنون نیز ادامه دارد، موجب پیدایش سامانه‌های ترجمه ماشینی متعددی شده است، که متون را به صورت خودکار از یک زبان مبدأ به یک زبان مقصد ترجمه می‌کند.

اما در حال حاضر که بیش از نیم قرن از پیدایش سامانه‌های ترجمه ماشینی می‌گذرد، این سامانه‌ها هنوز قادر به ارائه یک ترجمه مطلوب و عاری از اشتباه نیستند؛ هر چند که در این بازه زمانی بهبودهایی حاصل شده است.

عدم امکان جایگزینی مترجمان انسانی با سامانه‌های ترجمه ماشینی، محققان حوزه ترجمه ماشینی را بر آن داشته است تا سعی کنند با ارائه ابزارهایی، کار ترجمه را برای مترجمان تسهیل و تسریع نمایند؛ سامانه‌های مترجم‌یار تعاملی، آخرین نسخه از نتیجه این تلاش‌هاست.

ایده سامانه مترجم‌یار اولین بار در سال ۱۹۷۳ توسط مارتین کی (Kay, 1973) مطرح شد. در سامانه پیشنهادی او همکاری کاربر با ماشین به صورت پرسش و پاسخ‌هایی درباره معنی کلمات، تعیین مرجع ضمیر و موضوعاتی از این قبیل، انجام می‌شد. اما این نوع تعامل بین سامانه و کاربر مورد اقبال کاربران قرار نگرفت. از سال ۱۹۷۳ تا ۱۹۹۷ تعدادی از

محققان نظیر (Brown and Nirenburg, 1990; Maruyama and Watanabe, 1990; Whitelock et al., 1986), توسعه سامانه‌های مشابهی روی آوردند. در این سامانه‌ها سعی شد با تکنیک‌هایی نظیر مرتب‌سازی سؤالات به‌منظور کاهش سؤالات مورد نیاز و ارائه پاسخ‌های چندگزینه‌ای با انتخاب محتمل‌ترین جواب به‌عنوان انتخاب پیش‌فرض، رضایت کاربر جلب شود. هر چند که در سامانه‌های اخیر تلاش زیادی برای بهبود کارایی فرآیند رفع ابهام و کاهش نارضایتی کاربر در نحوه تعامل با سامانه انجام شد، اما همچنان فرآیند پرسش و پاسخ در این سامانه‌ها باقی ماند و جایگزین مناسبی برای نحوه تعامل با کاربر ایجاد نشد. از آنجایی که این نحوه تعامل، مورد رضایت مترجمان نبود، تنها در کاربردهایی که هزینه ترجمه دستی به میزان کافی بالا بود، کاربران به استفاده از این سامانه‌ها روی می‌آوردند. نظیر مواردی که کاربران دانش لازم را برای ترجمه نداشتند. در سال ۱۹۹۷ با معرفی پروژه ترن‌تایپ در (Foster et al., 1997)، یک تحول بزرگ در نحوه تعامل کاربر با ماشین، رخ داده شد. در سامانه ترن‌تایپ مشارکت ماشین در فرآیند ترجمه، در قالب پیشنهادهایی درباره جمله مقصد انجام می‌شد؛ درحالی‌که مترجم می‌توانست در روش‌های متنوعی در این تعامل، شرکت نماید. این روش‌ها شامل تایپ قسمت‌هایی از متن مقصد و انجام اصلاحات روی پیشنهادهای سامانه است. در این سامانه، مترجم بر تمام مراحل فرآیند ترجمه کنترل دارد و این سامانه است که باید درون محدودیت‌های ضمنی ایجاد شده از طریق مشارکت‌های مترجم کار کند. مترجم می‌تواند آزادانه پیشنهادهای ارائه شده توسط سامانه را بپذیرد، تغییر دهد و یا از آن‌ها صرفه نظر نماید. پس از ارائه فهرست ترجمه‌های پیشنهادی در بخش مشارکت سامانه و اعمال نظر توسط مترجم، بخش تولیدشده از متن هدف به‌عنوان یک بخش تأییدشده و صحیح، مبنای ترجمه قسمت باقیمانده قرار می‌گیرد و فهرست ترجمه‌های پیشنهادی در مشارکت بعدی سامانه باید منطبق با این بخش تأیید شده باشد. به این قطعه از ترجمه که مطابق با نظر کاربر تنظیم شده است در اصطلاح پیشوند^۱ گفته می‌شود. همچنین به ترجمه پیشنهادی سامانه برای تکمیل پیشوند، در اصطلاح پسوند^۲ گفته می‌شود.

از سال ۱۹۹۷ تا سال ۲۰۰۴ اغلب کارهای انجام شده و مقالات ارائه شده در زمینه سامانه‌های مترجم‌یار مربوط به نسخه‌های مختلف ترن‌تایپ هستند (Langlais et al., 2000 & Foster, 2002; Cubel et al., 2004) در سال ۲۰۰۵ یک روش جستجوی جدید برای ارائه پسوند در (Bender et al., 2005) ارائه شد. همچنین در (Civera et al., 2006; Barrachina et al., 2007) از ماشین حالات محدود برای ساخت گراف کلمه در سامانه مترجم‌یار تعاملی، استفاده شد. در سال ۲۰۰۹ نیز پروژه تحت وب کیترا در (Koehn, 2009a and 2009b) معرفی شد. همچنین در سال ۲۰۱۰ (Ortiz-Martínez et al., 2010)، یک سامانه مترجم‌یار تعاملی مجهز به یادگیری برخط با استفاده از بازخورد کاربر ارائه شد.

کاری که ما قصد داریم در این مقاله ارائه دهیم، از جهاتی مشابه با پروژه کیتراست. در این پروژه و کار ما، از سامانه ترجمه ماشینی موز (Koehn et al., 2007) به‌عنوان موتور ترجمه سامانه مترجم‌یار تعاملی استفاده می‌شود. اما برخلاف این تشابه، سامانه ما از جهاتی با این پروژه و سایر سامانه‌های مترجم‌یار نظیر سامانه ارائه‌شده در (Barrachina et al., 2007)، متفاوت است. در سامانه‌های یادشده از فاصله ویرایشی، برای یافتن منطبق‌ترین فرضیه ترجمه تولیدشده توسط موتور ترجمه، با پیشوند مورد تأیید کاربر استفاده می‌شود و اگر چندین فرضیه دارای فاصله ویرایشی حداقل باشند، کم‌هزینه‌ترین فرضیه انتخاب می‌شود. همان‌طور که ما در بخش سوم از این مقاله ارائه خواهیم داد، این روش به دلیل ضعف قابلیت‌های معیار فاصله ویرایشی، ایرادهایی دارد، که این موارد بر کارایی سامانه مترجم‌یار تأثیر مستقیمی دارد. ما در این مقاله با ارائه یک معیار کامل‌تر، سعی کرده‌ایم تا از تأثیرات منفی معیار فاصله ویرایشی بکاهیم و کارایی سامانه مترجم‌یار تعاملی را افزایش دهیم.

در ادامه این مقاله، در بخش دوم به معرفی موتور ترجمه سامانه مترجم‌یارمان، که درواقع یک سامانه ترجمه ماشینی آماری است، می‌پردازیم. در بخش سوم رویکرد تعاملی سامانه مترجم‌یارمان و روش‌های جستجوی پیشوند در گراف جستجو را بررسی می‌کنیم. همچنین در بخش چهارم، نتایج آزمایش‌های انجام‌شده روی سامانه مترجم‌یار تعاملی را ارائه خواهیم داد و درنهایت در بخش پنجم یک نتیجه‌گیری از کارهای انجام شده ارائه خواهد شد و روال‌های آتی این پروژه نیز معرفی می‌شود.

¹ prefix
² suffix

۲- موتور ترجمه

همان‌طور که در بخش مقدمه اشاره شد، در سامانه مترجم‌یار تعاملی که ما آن را پیاده‌سازی کرده‌ایم، از سامانه ترجمه ماشینی موز (Koehn et al., 2007) به‌عنوان موتور ترجمه استفاده شده است. موز یک سامانه ترجمه ماشینی آماری مبتنی بر عبارات است که این امکان را فراهم می‌آورد تا مدل‌های ترجمه برای هر جفت زبان دلخواه به‌صورت خودکار یادگرفته شوند. به‌طور کلی ما از موز در دو موقعیت استفاده می‌کنیم: (۱) از بهترین فرضیه ترجمه تولیدشده توسط موز، به‌عنوان یک ترجمه پیشنهادی کامل در شروع تعامل با کاربر استفاده می‌کنیم. (۲) از فرضیه‌های ترجمه تولیدشده در مرحله رمزگشایی موز، برای ساخت یک گراف جستجو، که در بخش تعاملی مورد استفاده قرار می‌گیرد، استفاده می‌کنیم.

برای معرفی بهتر موز لازم است تا به‌صورت مختصر با نحوه عملکرد یک سامانه ترجمه ماشینی آماری آشنا شویم. همچنین لازم است تا قدری به معرفی مرحله رمزگشایی موز بپردازیم تا با نحوه ساخت گراف جستجو که در تعامل با کاربر، نقش اساسی ایفا می‌کند، بهتر آشنا شویم. در ادامه به بررسی موارد یاد شده می‌پردازیم.

۲-۱- معرفی سامانه ترجمه ماشینی آماری

ترجمه ماشینی آماری فرآیندی است که طی آن به‌صورت خودکار یک جمله از زبان مبدأ با استفاده از پیکره‌های موازی دوزبانه و ثنوری‌های آماری به جمله‌ای در زبان مقصد ترجمه می‌شود. در ترجمه یک متن باید هر کدام از جملات زبان مبدأ f_1^I به یک جمله در زبان مقصد e_1^I ترجمه شود. جمله f_1^I شامل کلمات $f_1, \dots, f_j, \dots, f_r$ است و e_1^I شامل کلمات $e_1, \dots, e_i, \dots, e_j$ می‌شود. توصیف تناظر بین دو جمله از زبان مبدأ و مقصد از احتمال شرطی $Pr(e_1^I | f_1^I)$ استفاده می‌کنیم. بدین ترتیب ترجمه‌ای برای جمله f_1^I انتخاب خواهد شد، که دارای بیشترین احتمال باشد، بنابراین با حل مسئله زیر می‌توان بهترین ترجمه را برای جمله مورد نظر پیدا کرد.

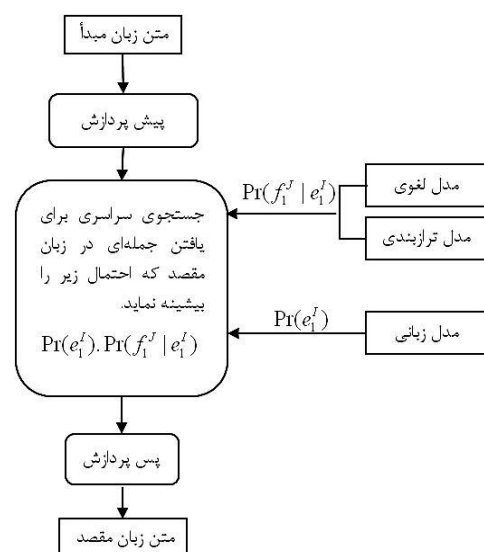
$$e_1^I = \operatorname{argmax}\{Pr(e_1^I | f_1^I)\} \quad (1)$$

$$= \operatorname{argmax}\{Pr(e_1^I).Pr(f_1^I | e_1^I)\} \quad (2)$$

احتمال $Pr(e_1^I | f_1^I)$ را با استفاده از قانون بیز به احتمالات $Pr(e_1^I)$ و $Pr(f_1^I | e_1^I)$ تجزیه می‌کنیم. با استفاده از این عمل

احتمال مورد نظر به دو مدل آماری تفکیک می‌شود. احتمال $Pr(e_1^I)$ مربوط به مدل زبانی است و مشخص می‌کند که به چه میزان جمله ترجمه‌شده از لحاظ شیوایی و درستی جمله، صحیح است. استفاده از مدل زبانی از ایجاد جملات نادرست در زبان مقصد جلوگیری می‌کند. احتمال $Pr(f_1^I | e_1^I)$ مربوط به مدل ترجمه است و بررسی می‌کند که با چه احتمالی جمله e_1^I می‌تواند ترجمه‌ای برای جمله مبدأ f_1^I باشد. در شکل (۱) معماری یک سامانه ترجمه ماشینی آن‌چنان که توضیح داده شد، دیده می‌شود. هدف از بررسی ترجمه ماشینی در این گزارش بررسی مدل ترجمه به‌کار رفته در آن است؛ لذا از تشریح مدل زبانی اجتناب کرده و به بررسی مدل‌های ترجمه می‌پردازیم.

مسئله مهم در مدل‌های ترجمه، آن است که چگونه تناظر بین اجزای جمله مبدأ و مقصد را در نظر بگیریم، به این تناظر، ترازبندی گفته می‌شود. براساس مدل‌های ترازبندی، مدل‌های ترجمه متفاوتی ایجاد شده است. در سامانه مترجم‌یار حاضر از مدل ترجمه مبتنی بر عبارات که مورد استفاده سامانه موز نیز می‌باشد، استفاده شده است.



(شکل ۱) - معماری سامانه ترجمه ماشینی مبتنی بر قانون تصمیم‌گیری بیز (Zens et al., 2002).

در مدل ترجمه مبتنی بر عبارات جهت انجام یک ترجمه صحیح و کامل از ترازبندی عبارات استفاده می‌شود، این مدل در (Zens et al., 2002) تشریح شده است. در این ایده که در ادامه به تفصیل تشریح می‌شود، برای ترجمه یک جمله از زبان مبدأ به زبان مقصد، ابتدا باید هر جمله زبان مبدأ را

به قطعاتی تقسیم کرد و سپس این قطعات ترجمه می‌شوند و ترکیب قطعات ترجمه‌شده، جمله مقصد را خواهند ساخت. هر یک از این قطعات مذکور را یک عبارت می‌نامیم. یک عبارت دنباله‌ای از کلمات متوالی در یک جمله است. در سامانه‌های مترجم‌یار آماری به‌طور عمومی از مدل ترجمه مبتنی بر عبارات استفاده می‌شود. چرا که با به‌کارگیری عبارات به‌عنوان واحد ترجمه، ترجمه دقیق‌تر و مرتبط‌تری با مضمون متن ورودی ایجاد می‌شود.

۲-۲- فاز رمزگشایی^۱ موزز

وظیفه فرآیند رمزگشایی یک سامانه ترجمه ماشینی، پیدا کردن بهترین ترجمه ماشینی بر طبق امتیازات به‌دست‌آمده از مدل ترجمه و مدل زبانی است. مسأله پیدا کردن بهترین ترجمه و یا به عبارت دیگر محتمل‌ترین ترجمه، یک مسأله دشوار است؛ چرا که تعداد ترجمه‌های ممکن برای یک جمله ورودی به‌صورت نمایی با طول جمله، رشد می‌کند و نشان داده شده است که مسأله رمزگشایی، ترجمه ماشینی یک مسأله NP-complete است (knight, 1999). بنابراین بررسی جامع همه ترجمه‌های ممکن یک جمله، امتیازدهی آنها و انتخاب بهترین ترجمه از لحاظ محاسباتی بسیار پرهزینه است و درعمل حتی برای جملاتی با طول کوتاه و متوسط نیز، این روش غیرکاربردی است. لذا برای حل این مسأله از روش‌های جستجوی اکتشافی^۲ استفاده می‌شود. برای این منظور، در موزز از روش جستجوی شعاعی^۳ که یک روش جستجوی اکتشافی است، استفاده شده است. در روش جستجوی شعاعی، ترجمه یک جمله با اضافه‌کردن ترجمه عبارات، کامل می‌شود. در طول جستجو، ترجمه‌های جزئی ساخته می‌شوند، که آنها را فرضیه می‌نامیم. از نقطه‌نظر برنامه‌نویسی، یک فرضیه یک ساختمان داده است که اطلاعاتی را درباره ترجمه جزئی دربر می‌گیرد. برای نمونه، اطلاعات راجع به اینکه چه کلماتی از جمله ورودی تاکنون ترجمه شده‌اند؛ چه کلماتی در خروجی تولید شده‌اند؛ امتیاز احتمالاتی ترجمه جزئی چقدر بوده است و مواردی از این قبیل.

برای ترجمه یک جمله، اولین فرضیه‌ای که مورد استفاده قرار می‌گیرد، یک فرضیه تهی است. این فرضیه هیچ کلمه‌ای از جمله مبدأ را پوشش نمی‌دهد و ترجمه‌ای را نیز در خروجی تولید نمی‌کند و امتیاز احتمالاتی جزئی این

فرضیه به‌طور قراردادی، یک درنظر گرفته می‌شود. با اضافه‌شدن ترجمه یکی از عبارات مبدأ، فرضیه تهی توسعه داده شده و یک فرضیه جدید ایجاد می‌شود. در توسعه‌های بعدی این فرضیه، عبارت ترجمه‌شده از جمله مبدأ دیگر نباید دوباره ترجمه شود. همچنین پس از توسعه هر فرضیه، همه هزینه‌های احتمالاتی ترجمه جزئی جدید محاسبه می‌شود؛ این هزینه‌ها شامل هزینه ترجمه عبارت مبدأ پوشش داده شده اخیر و هزینه تخمینی توسعه‌های بعدی فرضیه تا مبدل شدن به یک فرضیه کامل، است.

فرآیند توسعه فرضیه‌های جزئی که در پاراگراف قبل شرح داده شد، تا زمانی که همه فرضیه‌ها به‌طور کامل توسعه داده شوند و به یک فرضیه کامل تبدیل شوند، به‌طوری‌که همه عبارات جمله مبدأ را پوشش دهد، ادامه می‌یابد.

سامانه ترجمه ماشینی موزز برای اعمال روش جستجوی شعاعی، از هرس فرضیه‌های جزئی بدتر، بهره می‌برد؛ بدین ترتیب فرضیه‌هایی که تعداد کلمات یکسانی را از جمله مبدأ تحت پوشش قرار می‌دهند، با هم مقایسه می‌شوند و فرضیه‌هایی که مجموع هزینه‌های جاری و آتی آنها بیشتر از یک آستانه از پیش مشخص شده باشد، هرس می‌شوند. در روش دیگری از هرس، برای هر دسته از فرضیه‌ها که کلمات پوشش داده‌شده یکسانی دارند، یک تعداد بیشینه در نظر گرفته می‌شود و فرضیه‌های اضافی که هزینه‌های بیشتری را به سامانه تحمیل می‌کنند، دور ریخته می‌شوند.

در مرحله رمزگشایی موزز، از ساختمان داده پشته به‌منظور یک ساختار مناسب برای هرس فرضیه‌های ترجمه استفاده می‌شود؛ به‌طوری‌که هر پشته تنها فرضیه‌هایی را نگهداری می‌کند که تعداد کلمات خارجی ترجمه‌شده آنها با هم برابر باشد؛ به‌عنوان مثال یک پشته فقط شامل فرضیه‌هایی است که یک کلمه خارجی را ترجمه کرده‌اند و پشته دیگر شامل فرضیه‌هایی با دو کلمه خارجی تحت پوشش است و به همین ترتیب پشته‌های متعددی تعریف می‌شوند.

در شکل (۲) نمونه‌ای از سازماندهی فرضیه‌ها درون پشته‌ها نشان داده شده است. پشته شماره صفر، همواره یک فرضیه دارد که آن هم فرضیه تهی است. از طریق توسعه فرضیه‌ها مطابق با روشی که درقبل شرح داده شد، فرضیه‌های جدید ساخته شده و درون پشته مربوطه خود قرار می‌گیرند. فرضیه‌های ترجمه موجود در شکل (۲)، برای

¹ Decoding

² Heuristic

³ Beam Search

موجود نباشد، سامانه را قادر کند که ترجمه قابل قبولی به کاربر ارائه دهد.

۳-۱- روش جستجوی مبتنی بر فاصله ویرایشی

مطابق با (Koehn et al., 2007; Barrachina et al., 2007)، برای ارائه یک پیشنهاد مکمل با پیشنهاد به کاربر، ابتدا می‌بایست گره‌ای در گراف جستجو، پیدا کنیم که با عبارات پیشنهادی که در تعاملات قبلی کاربر با سامانه به وجود آمده است، کمترین فاصله ویرایشی را داشته باشد؛ ما نام این روش را جستجوی مبتنی بر فاصله ویرایشی می‌نامیم. منظور از فاصله ویرایشی بین دو رشته از کلمات، حداقل تعداد لازم عملیات درج، حذف و جایگزینی است تا یک رشته به‌طور کامل مطابق با رشته دیگر شود (Levenshtein, 1965).

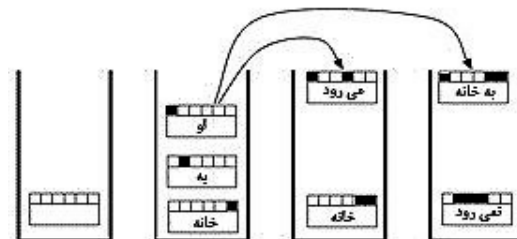
این روش بر پایه این فرض بنا شده است که، فرضیه جزئی که کمترین فاصله ویرایشی را با پیشنهاد داشته شانس بیشتری دارد که در ادامه مسیر تا تبدیل شدن به یک فرضیه کامل، با ترجمه پسوند مورد نظر کاربر منطبق باشد.

یک مثال از این روش در شکل (۳) دیده می‌شود. در این مثال، ما قصد داریم، جمله فارسی "ایده مرد جوان بیش از حد مخاطره‌آمیز است"، را با سامانه مترجم‌یار به انگلیسی ترجمه کنیم. فرض می‌کنیم جمله انگلیسی مورد نظر مترجم، "the idea of the young man is too risky" و پیشنهاد حاضر نیز "the idea o" باشد. در گراف جستجوی شکل (۳)، سه فرضیه با حداقل فاصله ویرایشی یک وجود دارد، لذا برای انتخاب منطبق‌ترین فرضیه می‌بایست هزینه ترجمه فرضیه‌ها مورد مقایسه قرار گیرد.

منظور از هزینه ترجمه، مجموع هزینه ترجمه فرضیه جاری و هزینه تخمینی که این فرضیه تا تبدیل شدن به یک فرضیه ترجمه کامل متحمل می‌شود، است. در نهایت فرضیه چهار به‌عنوان کم‌هزینه‌ترین فرضیه منطبق با پیشنهاد انتخاب شده و کم‌هزینه‌ترین مسیر مکمل این فرضیه به کاربر پیشنهاد داده می‌شود. بر اساس فرضیه انتخاب شده، پیشنهاد سامانه برای تکمیل پیشنهاد مذکور پسوند "f the young man is too risky" است.

ترجمه جمله انگلیسی "He does not go to home" به یک جمله فارسی تشکیل شده‌اند.

در بخش بعدی این مقاله، ما به ساختار تعاملی سامانه مترجم‌یار طراحی شده می‌پردازیم.



(شکل ۲) - سازماندهی فرضیه‌ها در ساختمان داده پشته. تعداد پشته‌ها همواره یکی بیشتر از تعداد کلمات جمله مبدأ است. پشته اضافی مربوط به فرضیه تهی است. فرضیه‌های درون پشته‌ها با اشاره‌گرهایی به هم مرتبط هستند (Koehn et al., 2010).

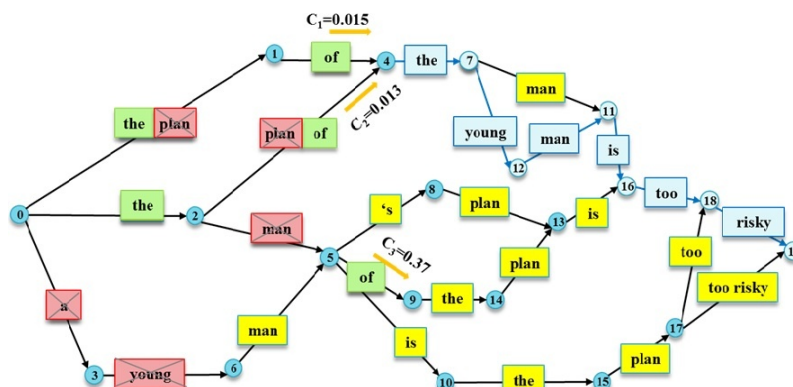
۳-۲- موتور تعاملی

همان‌طور که در بخش مقدمه شرح داده شد، پس از هر دفعه‌ای که کاربر با استفاده از دکمه‌های صفحه کلید تغییری را در ترجمه ارائه شده توسط سامانه، اعمال می‌کند، سامانه بلافاصله با توجه به ترجمه اصلاح شده، ترجمه تکمیل‌کننده‌ای را ارائه می‌دهد.

حال در این بخش می‌خواهیم بررسی کنیم که چگونه سامانه قادر است عبارات تکمیلی را بر اساس ترجمه پیشنهادی ارائه کند. برای این منظور لازم است تا یک گراف جستجو با استفاده از فرضیه‌های ترجمه جزئی، که در مرحله رمزگشایی موزز تولید شده‌اند، ساخته شود و در مرحله بعدی، این گراف برای یافتن یک ترجمه منطبق با قطعه ترجمه مورد تأیید کاربر، مورد جستجو و پویش قرار گیرد.

در واقع وظیفه اصلی سامانه تعاملی آن است که، با استفاده از ترجمه‌های ممکن در گراف جستجو و با توجه به پیشنهادی که کاربر، ترجمه تکمیل‌کننده‌ای به کاربر پیشنهاد دهد.

در ادامه، نخست معیارهای مناسب برای جستجوی پیشنهادی در گراف جستجو را بررسی می‌کنیم و به ارائه روش پیشنهادی‌مان در این زمینه می‌پردازیم؛ سپس به بررسی روش‌هایی می‌پردازیم تا در مواقعی که ترجمه مورد نظر کاربر، بنا به دلایلی نظیر کامل نبودن مدل‌های آماری و یا هرس فرضیه‌ها در مرحله رمزگشایی، در گراف مورد جستجو



(شکل ۳) - جستجوی پیشوند "the idea o" در گراف جستجوی مبتنی بر عبارات، با استفاده از روش مبتنی بر فاصله ویرایشی.

۳-۱- روش مبتنی بر مجموع وزن دار فاصله

ویرایشی و هزینه ترجمه

روش جستجوی مبتنی بر فاصله ویرایشی یک ایراد اساسی دارد که مربوط به استفاده از معیار فاصله ویرایشی جهت سنجش میزان تطابق فرهنگیها پیشوند می شود. در محاسبه فاصله تنها عملیات درج، حذف و جایگزینی در نظر گرفته می شود و به جابه جایی عبارات توجهی نمی شود.

برای روشن تر ساختن این مشکل از یک مثال استفاده می کنیم. فرض می کنیم در ترجمه انگلیسی به فارسی، جمله مورد نظر کاربر جمله فارسی "مولانا بهترین شاعر قرن هفتم، یکی از مشاهیر بزرگ جهان است که قالب شعرهای او مثنوی است" باشد. همچنین دو فرضیه ترجمه در دسترس باشد. فرضیه اول با هزینه ترجمه ۰/۱۲، جمله فارسی "یکی از مشاهیر بزرگ جهان، مولانا بهترین شاعر قرن هفتم است که قالب شعرهای او مثنوی می باشد"، است و فرضیه دوم با هزینه ترجمه ۰/۶۸۱ به صورت "قالب مولانا بهترین شاعر قرن هفتم، یکی از مثنوی بزرگ شعر است که جهان مشاهیر است"، می باشد.

اگر پیشوند ما برابر با رشته "مولانا بهترین شاعر قرن هفتم، یکی از مشاهیر بزرگ جهان است ک" باشد، بنابراین فاصله ویرایشی فرضیه اول، برابر با ۱۰ خواهد شد. این در حالی است که فاصله ویرایشی فرضیه دوم، که از لحاظ ساختار و معنا جمله صحیحی نیست، سه است.

طبق روش مبتنی بر فاصله ویرایشی فرضیه دوم به عنوان فرضیه منطبق با پیشوند انتخاب شده و بر این اساس پسوند "ه جهان مشاهیر است" به کاربر پیشنهاد داده می شود که مورد تأیید کاربر قرار نخواهد گرفت.

همان طور که در مثال قبل دیدیم، اگر ما تنها روی فاصله ویرایشی فرضیه با پیشوند تکیه کنیم، ممکن است جستجوی ما منجر به فرضیه ای شود که از لحاظ مدل های ترجمه و زبانی، امتیاز احتمالاتی پایینی داشته باشد و طبیعتاً چنین فرضیه ای مطلوب کاربر نخواهد بود. هر چند که فاصله ویرایشی آن با پیشوند حداقل باشد.

برای حل این مشکل ما استفاده از مجموع وزن دار هزینه ترجمه و فاصله ویرایشی را پیشنهاد می دهیم. ایده این روش از این حقیقت ناشی می شود که ترجمه ای که از لحاظ مدل های آماری زبانی و ترجمه، امتیاز بهتری را کسب می کند، احتمال بیشتری دارد که یک ترجمه صحیح باشد و با ترجمه نهایی مورد نظر کاربر تطابق داشته باشد. هر چند که این ترجمه ممکن است در مقاطعی با پیشوندهای ترجمه کاربر به دلیل مسأله جابه جایی عبارات، کمترین فاصله ویرایشی را نداشته باشد. شبهه کد مربوط به روش جستجوی پیشنهادی در شکل (۴) آمده است.

بر طبق این الگوریتم، در طی فزآید جستجو ی پیشوند در گراف جستجو، ما از گره اول گراف (فرضیه تهی) شروع کرده و گره های مختلف را بررسی می کنیم تا به گره ی برسیم که مجموع وزن دار هزینه های ترجمه و فاصله ویرایشی آن تا پیشوند از سایر گره های گراف کمتر باشد؛ با رسیدن به هر گره ما یک یا چندین اگوارمه عقب برای آن گره تعریف می کنیم.

ترجمه مورد تأیید کاربر نزدیک است و هم از لحاظ مدل‌های آماری ترجمه و مدل زبانی، ترجمه صحیح‌تری است؛ چنین فرضیه‌ای می‌تواند نویدبخش ارائه ترجمه مکمل بهتری به کاربر باشد.

همچنین لازم به ذکر است که الگوریتم ارائه‌شده هیچ محدودیتی برای جملات خاص زبان، مانند جملات برگشتی^۱ ندارد؛ البته این ادعا در شرایطی صدق می‌کند که موتور ترجمه استفاده‌شده در سامانه مترجم‌یار محدودیتی برای ترجمه این جملات نداشته باشد.

وزن مربوط به فاصله ویرایشی و هزینه ترجمه در آزمایش‌ها با استفاده از پیکره توسعه و به‌صورت تجربی تعیین می‌شود؛ در آزمایش‌های ارائه‌شده در بخش چهارم مقاله، این وزن برای فاصله ویرایشی ۰/۲ و برای هزینه ترجمه ۰/۸ به‌دست آورده شده است.

۳-۲- استفاده از روش‌های Back-off

در مواردی ممکن است، هیچ یک از روش‌های جستجوی شرح داده شده در دو بخش قبل، قادر به ارائه یک پیشنهاد به کاربر نباشند. این مشکل در مواردی پیش می‌آید که جزء آخر پیشوند یک کلمه ناتمام باشد و فرضیه‌ای هم در گراف جستجو موجود نباشد که شامل کلمه منطبق با این کلمه ناتمام باشد.

برای مثال اگر جزء آخر پیشوند کلمه ناتمام "آر" باشد، تنها فرضیه‌هایی که کلماتی مانند "آرمان، آرش، آرام و ..." دارند، می‌توانند به‌عنوان یک فرضیه منطبق با پیشوند در نظر گرفته شوند. اگر هیچ فرضیه‌ای در گراف وجود نداشته باشد که دارای کلمه‌ای باشد که با کلمه ناتمام "آر" منطبق باشد، هر پیشنهادی از سوی سامانه براساس گراف جستجو، مورد تأیید کاربر نخواهد بود.

در چنین مواردی ما نیاز داریم تا از روش‌های دیگر برای ارائه یک پیشنهاد مناسب استفاده کنیم، هر چند که ممکن است این پیشنهاد قطعیت کمتری داشته باشد. چندین راه حل برای حل این مسأله تاکنون پیشنهاد شده است. در (Barrachina et al., 2007) استفاده از مدل زبانی، برای یافتن یک عبارت تکمیلی با بالاترین احتمال پیشنهاد شده است. در این روش، دیگر از مدل‌های ترجمه ماشینی استفاده نمی‌شود و تنها کلمه یا عبارتی به‌عنوان پسوند انتخاب می‌شود که بالاترین احتمال را بر حسب مدل زبانی کسب می‌کند. برای این منظور، ابتدا می‌بایست از فهرست

^۱ Recursive

```

Input: user_prefix  $u$ , search_graph  $g$ ,
weight_translation_cost  $\alpha$ 
Output: Best_Adapt_Hypothesis
1. allowable_cost & error total_cost= MAX_INT
2. Adapt_Hypothesis= $\emptyset$ 
3. add backpointer ( cost=0.0, error=0, toProcess= $u$ 
, Father_Hyp=NULL) to start node
4. for all node  $s \in g$  in topologically increasing order do
5.   for all backpointer  $b$  of node  $s$  do
6.     if  $\alpha \times b.cost + (1 - \alpha) \times b.error \leq total\_cost$  then
7.       for all transition  $t$  from node  $s$  do
8.         compute string edit distance matrix for
            $b.toProcess, t.phrase$ 
9.         for all matches  $m$  in matrix that consumed all of
            $b.toProcess$  do
10.          new cost  $c' = s.cost + t.cost$ 
           +  $t.toState.forwardCost$ 
11.          new error  $e' = s.error + m.error$ 
12.          if  $\alpha \times c' + (1 - \alpha) \times e' < total\_cost$  then
13.            set this as Adapt_Hypothesis
14.             $total\_cost = \alpha \times c' + (1 - \alpha) \times e'$ 
15.          end if
16.        end for
17.      for all matches  $m$  in the matrix that consumed
           all of  $t.phrase$  do
18.        reached new state  $s_n = t.toState$ 
19.        create new backpointer  $b_n$ 
20.         $b_n.cost = s.cost + t.cost$ 
21.         $b_n.error = s.error + m.error$ 
22.         $b_n.toProcess = s.toProcess - t.phrase$ 
23.         $b_c$  = current backpointer for state  $s_n$  at prefix
           pos.  $b_n.toProcess$ 
24.        if  $b_c$  not defined or  $(\alpha \times b_n.cost + (1 - \alpha) \times b_n.error < \alpha \times b_c.cost + (1 - \alpha) \times b_c.error)$  then
25.          make  $b_n$  new backpointer for state  $s_n$  at prefix
           pos.  $b_n.toProcess$ 
26.        end if
27.      end for
28.    end for
29.  end if
30. end for
31. end for
32. Best_Adapt_Hypothesis = Adapt_Hypothesis

```

(شکل ۴) - الگوریتم جستجوی پیشوند مبتنی بر مجموع وزن‌دار فاصله ویرایشی و هزینه ترجمه

از آنجا که برای رسیدن به یک گره از گراف مسیرهای متعددی وجود دارد، همچنین هر مسیر خطای فاصله ویرایشی و هزینه ترجمه مخصوص به خود را دارد، می‌بایست اشاره‌گری را برای مشخص کردن فرضیه والد و خطای ویرایشی فرضیه جاری و هزینه ترجمه آن، تعریف کرد.

علاوه بر این ما با استفاده از اشاره‌گر می‌توانیم مشخص کنیم که فرضیه جاری تاکنون، چه بخشی از پیشوند را پوشش داده و چه بخشی از پیشوند باقیمانده است و باید با فرضیه‌های توسعه یافته بعدی منطبق شود. در نهایت با پویای گراف جستجو فرضیه‌ای را می‌یابیم که هم به

۱- گرام‌های مدل زبانی، تمام کلماتی را که با کلمه ناتمام انتهای پیشوند، آغاز می‌شوند، استخراج کرد و سپس احتمال n -گرام تک تک کلمات استخراج شده را به ازای $n-1$ کلمه ماقبل آخر پیشوند محاسبه کرد. هر کلمه‌ای که احتمال رخداد آن به ازای کلمات ماقبل آخر پیشوند بیشتر باشد، به‌عنوان یک پیشنهاد جزئی به کاربر ارائه خواهد شد. این روش یک ایراد عمده دارد، و آن این است که جمله مبدا را اصلاً در نظر نمی‌گیرد.

برای مثال فرض می‌کنیم، جمله مبدا ما، جمله انگلیسی "the sound of the music was loud" و پیشوند مورد بررسی "صدای آه" باشد. در این مثال مدل زبانی ۲-گرام، از میان کلماتی مانند "آه، آهسته، آهو، آهنگ، آهن و آهک"، کلمه "آهسته" را بر اساس احتمال رخداد آن پس از کلمه "صدای" به‌عنوان یک پیشنهاد به کاربر، انتخاب می‌کند؛ در حالی که اگر به جمله زبان مبدا نیز توجهی شده بود، این احتمال وجود داشت که کلمه "آهنگ" انتخاب شود.

برای حل این مشکل، از مجموع وزن دار مدل لغوی IBM-1 (Brown et al., 1993) و مدل زبانی استفاده می‌کنیم. دلیل انتخاب مدل IBM-1، آن است که مدل‌های بالاتر IBM قیود بیشتری نسبت به مدل IBM-1 دارند؛ و این محدودیت‌های بیشتر موجب می‌شود در شرایط نامناسبی که سامانه قادر به ارائه یک ترجمه مکمل نیست، مدل ترجمه شانس کمتری برای یافتن یک ترجمه مکمل داشته باشد.

در مدل IBM-1، برای به‌دست آوردن احتمال اینکه جمله زبان مقصد $E=e_i^l$ ترجمه جمله زبان مبدا $F=f_1^l$ باشد، به‌صورت زیر عمل می‌شود:

$$\Pr(E|F) = \prod_{i=1}^l \frac{1}{J+1} \sum_{j=1}^J p(e_i|f_j) \quad (3)$$

با توجه به معادله بالا، اگر بخواهیم احتمال اینکه کلمه e جزئی از ترجمه جمله f_1^l باشد، را محاسبه کنیم، باید به‌صورت زیر عمل کنیم:

$$\Pr(e|f_1^l) = \frac{1}{J+1} \sum_{j=1}^J p(e|f_j) \quad (4)$$

با توجه به توضیحات ارائه شده، روش کار بدین صورت خواهد شد، که ابتدا مانند روش قبل از فهرست ۱-گرام‌های مدل زبانی، تمام کلماتی را که با کلمه ناتمام انتهای پیشوند، آغاز می‌شوند، استخراج می‌کنیم. سپس با استفاده از مدل IBM-1 و طبق فرمول (۱۳)، احتمال آن را که هر یک از

کلمات استخراج‌شده، مربوط به ترجمه جمله ورودی باشند، محاسبه می‌کنیم. حال می‌بایست برای هر یک از احتمالات مدل IBM-1 و مدل زبانی، وزنی به‌صورت زیر در نظر گیریم.

$$\alpha P_{IBM1}(e_i|f_1^l) + (1-\alpha) P_{LM}(e_i|e_{i-1}) \quad (5)$$

در معادله بالا، مدل زبانی ۲-گرام در نظر گرفته شده است. همچنین، وزن α در آزمایش‌ها با استفاده از پیکره توسعه و به‌صورت تجربی تعیین می‌شود؛ در آزمایش‌های ارائه‌شده در بخش چهار این مقاله، این وزن ۰/۸۵ تعیین شده است.

۴- ارزیابی روش‌های پیشنهادی

در این بخش قصد داریم، سامانه مترجم‌یار تعاملی پیاده‌سازی شده بر اساس روش‌های پیشنهادی مان را در مقایسه با روش‌های قبلی ارزیابی کنیم؛ اما قبل از ارائه نتایج آزمایش‌های انجام شده و بررسی آنها، لازم است توضیحاتی در رابطه با نحوه شبیه‌سازی کاربر، معیارهای ارزیابی و پیکره‌های آموزش، توسعه و آزمون، ارائه دهیم؛ در ادامه به تشریح تمامی این موارد می‌پردازیم.

۴-۱- شبیه‌سازی مترجم انسانی

هدفی که به‌طور عمومی در سامانه‌های مترجم‌یار دنبال می‌شود، صرفه‌جویی در زمان کاربر (حرفه‌ای) در ترجمه متون است؛ که به‌عبارت دیگر، هدف بالابردن کارایی مترجم‌های انسانی است؛ ولی از آنجایی که استفاده از انسان برای ارزیابی سامانه، کار بسیار زمان‌بری است و مشکلاتی از قبیل یادگیری انسان در طول فرآیند ترجمه و در نتیجه تکرارناپذیر بودن آن و عدم یک‌پارچگی میان آزمایش‌های مختلف را دارد، امکان‌پذیر نیست (Megerdounian and Khadivi, 2010). بنابراین جهت ارزیابی سامانه‌های مترجم‌یار به‌طور عمومی از کاربران شبیه‌سازی شده استفاده می‌شود.

به‌منظور شبیه‌سازی عملکرد یک مترجم انسانی، ترجمه‌هایی که یک کاربر حقیقی باید در ذهنش داشته باشد، با مراجع زبان مقصد متون مورد ترجمه، شبیه‌سازی می‌شوند. سپس برای هر جمله زبان مبدا، اولین ترجمه پیشنهادی سامانه با ترجمه مرجع متناظرش مقایسه می‌شود و طولانی‌ترین پیشوند مشترک بین ترجمه پیشنهادی و ترجمه مرجع مشخص می‌شود. اولین نویسه غیر منطبق از ترجمه پیشنهادی باید با نویسه مرجع متناظرش جایگزین

استفاده شده است. اطلاعات آماری این پیکره انگلیسی-فارسی در جدول (۱) آمده است.

این پیکره به سه قسمت مجزا و بدون هم‌پوشانی تقسیم شده است و بدین ترتیب سه مجموعه آموزش، توسعه و ارزیابی را تشکیل داده است.

مجموعه آموزش و ارزیابی به ترتیب برای آموزش سامانه و ارزیابی آن مورد استفاده قرار می‌گیرند. از مجموعه توسعه نیز برای تنظیم پارامترهای روش‌های مورد ارزیابی، استفاده می‌شود.

(جدول ۱) - اطلاعات آماری پیکره انگلیسی-فارسی و ریموبیل

فارسی	انگلیسی	
۲۲۶۴۲		تعداد جملات
۲۳۳۹۴۸	۲۵۴۶۶۵	تعداد کلمات
۵۴۰۵	۲۶۹۶	واژگان
۲۵۰۱	۱۰۱۶	کلمات یک‌بار تکرار
۲۷۶		جملات
۳۳۳۹	۵۳۵۸	تعداد کلمات
۴۶۳	۵۱۱	واژگان
۲۰۰	۱۹۸	کلمات یک‌بار تکرار
۲۵۰		جملات
۲۶۹۲	۲۸۷۱	تعداد کلمات
۴۲۹	۳۴۵	واژگان
۱۹۳	۱۴۲	کلمات یک‌بار تکرار

۴-۴- بررسی نتایج آزمایش‌ها

اولین آزمایشی که در این بخش ارائه می‌شود، مربوط به ارزیابی روش جستجوی جدید پیشنهادی است. همان‌طور که در بخش ۲-۳ شرح دادیم، در روش جدیدی که ما ارائه دادیم، جستجوی پیشنهادی براساس مجموع وزن‌دار هزینه ترجمه و فاصله ویرایشی، انجام می‌شود، برخلاف روش مرسوم که تنها از فاصله ویرایشی جهت جستجوی پیشنهادی استفاده می‌شود. نتایج این آزمایش در جدول (۲) دیده می‌شود.

همان‌طور که از نتایج جدول (۲) مشخص است، استفاده از روش جستجوی مبتنی بر مجموع وزن‌دار فاصله ویرایشی و هزینه ترجمه، باعث کاهش معیار KSR و KSMR شده است. علاوه بر این معیارها، برای مقایسه دو روش، زمان لازم برای جستجوی پیشنهادی و ارائه یک پیشنهاد به کاربر، نیز باید مورد توجه قرار گیرد. طبق نتایج به‌دست آمده در جدول (۲)، به‌طور متوسط جستجو برای هر پیشنهاد در هر دو روش در کسری از ثانیه انجام می‌پذیرد.

شود. پس از این جابه‌جایی سامانه دوباره وارد عمل شده و پیشنهادهای جدیدی را ارائه می‌دهد. این فرآیند تا زمانی که ترجمه تولیدشده با ترجمه مرجع به‌طور کامل منطبق شود، ادامه می‌یابد.

کاربر برای یافتن طولانی‌ترین پیشنهادی که ترجمه پیشنهادی، درواقع به‌دنبال اولین خطای رخ داده شده، می‌گردد و با حرکت اشاره‌گر موقعیت پیشنهادی مورد نظرش را مشخص می‌کند. همچنین هر نویسه تایپ‌شده توسط کاربر نیز متناظر با یک ضربه به صفحه کلید است. اگر اولین نویسه نامنطبق، اولین نویسه ترجمه پیشنهادی باشد، در واقع هیچ پیشنهادی محاسبه نمی‌شود و اشاره‌گر نیز حرکت نخواهد کرد.

۴-۲- معیار ارزیابی

جهت ارزیابی سامانه مترجم‌بار از معیار نرخ ضربات صفحه کلید و عملکرد موشواره استفاده می‌شود. همان‌طور که در ادامه شرح خواهیم داد، این معیار به‌صورت مجموع تعداد ضربات صفحه کلید و تعداد حرکات موشواره تقسیم بر تعداد کل نویسه‌های مرجع محاسبه می‌شود. (Barrachina et al., 2007)

- نرخ ضربات به صفحه کلید (KSR)

$$KSR = \frac{\text{تعداد ضربات به صفحه کلید}}{\text{تعداد کل کاراکترهای مرجع}}$$

- نرخ حرکت اشاره‌گر موشواره (MAR)

$$MAR = \frac{\text{حرکات تعداد موشواره}}{\text{مرجع کاراکترهای کل تعداد}}$$

- نرخ ضربات به صفحه کلید و حرکت اشاره‌گر موشواره (KSMR)

$$KSMR = \frac{\text{تعداد حرکات موشواره} + \text{تعداد ضربات به صفحه کلید}}{\text{تعداد کل کاراکترهای مرجع}}$$

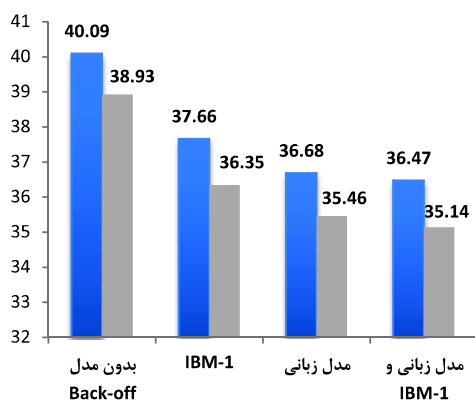
در بخش ۴-۴، از معیارهای فوق برای سنجش کارایی سامانه مترجم‌بار تعاملی استفاده خواهیم کرد.

۴-۳- معرفی پیکره آموزش، توسعه و ارزیابی

برای ارزیابی سامانه مترجم‌بار طراحی‌شده، از بخشی از پیکره پروژه و ریموبیل دانشگاه آخن (Ney et al., 2010)

(جدول ۳) - نتایج مربوط به ارزیابی روش‌های Back-off روی مجموعه ارزیابی پیکره ورموبیل، با استفاده از روش مبتنی بر فاصله ویرایشی برای جستجوی پیشوند در گراف.

روش جستجو مبتنی بر	روش Back-off	KSR (%)	KSMR (%)
فاصله ویرایشی	IBM-1	۲۵/۴۷	۳۷/۶۶
	مدل زبانی	۲۴/۳۹	۳۶/۶۸
	مدل زبانی + IBM-1	۲۴/۱۸	۳۶/۴۷
مجموع وزن دار فاصله ویرایشی و هزینه ترجمه	IBM-1	۲۴/۶۴	۳۶/۳۵
	مدل زبانی	۲۳/۵۹	۳۵/۴۶
	مدل زبانی + IBM-1	۲۳/۳۱	۳۵/۱۴



■ جستجوی مبتنی بر فاصله ویرایشی

■ جستجوی مبتنی بر مجموع وزن دار فاصله ویرایشی و هزینه ترجمه

(شکل ۵) - نمودار مربوط به نتایج حاصل از روش‌های جستجوی مبتنی بر فاصله ویرایشی و جستجوی مبتنی بر مجموع وزن دار هزینه ترجمه و فاصله ویرایشی

لازم به ذکر است که، مدل زبانی استفاده شده در این آزمایش‌ها یک مدل زبانی ۲-گرام است که روی مجموعه آموزش پیکره ورموبیل، آموزش دیده است. همان‌طور که از نتایج مشخص است و انتظار نیز داشتیم، ترکیب وزن دار مدل زبانی و مدل ترجمه IBM-1 بهترین نتایج را به وجود آورده‌اند. دلیل این نتیجه نیز مشخص است؛ اگر چه مدل IBM-1 این توانایی را دارد تا با قطعیت بیشتری نسبت به مدل زبانی، از میان کلمات ممکن مکمل با جزء ناتمام پیشوند، کلمه صحیح را انتخاب نماید،

اگر چه این زمان برای روش مبتنی بر فاصله ویرایشی قدری کمتر است، ولی در عوض روش مبتنی بر مجموع وزن دار فاصله ویرایشی و هزینه ترجمه، با ارائه پیشنهاد‌های صحیح‌تر، تعداد تعاملات انجام شده با کاربر را به میزان ۱۲۶ واحد کاهش داده است. همان‌طور که انتظار داشتیم، روش جستجوی پیشنهادی ما نتایج بهتری را کسب کرد. طبق نظریه‌ای که در بخش ۲-۳ ارائه دادیم، فرضیه‌هایی که از لحاظ مدل‌های ترجمه و مدل زبانی، امتیاز بهتری را کسب کرده‌اند باید شانس بیشتری برای تطابق با ترجمه مورد نظر کاربر داشته باشند.

فاصله ویرایشی، معیار کاملی برای بررسی میزان تطابق فرضیه‌های ترجمه با ترجمه مورد نظر کاربر نیست؛ چرا که امکان جابه‌جایی مجاز عناصر جمله را در نظر نمی‌گیرد و چنین جابه‌جایی‌هایی را در جمله، به‌عنوان فاصله ویرایشی لحاظ می‌کند.

در روش جستجوی پیشنهادی سعی شده است با استفاده از معیار هزینه ترجمه، از شدت تأثیر معیار فاصله ویرایشی کاسته شده و به فرضیه‌هایی که هزینه ترجمه کمتری دارند، توجه بیشتری شود.

آزمایش بعدی مربوط به استفاده از روش‌های Back-off است. روش‌های Back-off در واقع مکمل روش‌های جستجوی گراف بودند و در مواقعی که هیچ فرضیه‌ای در گراف یافت نشود که با پیشوند مطابق باشد، از روش‌های Back-off برای ارائه یک پیشنهاد، البته با قطعیت کمتر، استفاده می‌شد. در این آزمایش کارایی روش‌های مختلف Back-off با هم مقایسه می‌شود؛ نتایج این آزمایش در جدول (۳) قابل مشاهده است. همچنین در نمودار شکل (۵) نتایج هر دو آزمایش اخیر جهت تحلیل بهتر نشان داده شده است.

(جدول ۲) - نتایج مربوط به ارزیابی روش‌های جستجوی پیشوند روی مجموعه ارزیابی پیکره ورموبیل.

روش جستجو مبتنی بر	تعداد تعامل با کاربر	زمان متوسط (ثانیه)	KSR (%)	KSMR (%)
فاصله ویرایشی	۵۸۹۶	۰/۴۹	۲۹/۲۷	۴۰/۰۹
مجموع وزن دار فاصله ویرایشی و هزینه ترجمه	۵۷۷۰	۰/۶۶	۲۸/۵۸	۳۸/۹۳

مولوی است"، دو فاصله ویرایشی با هم دارند. در صورتی که اگر علاوه بر عملیات درج، حذف و جایگزینی از پرش عبارات، نیز استفاده می‌شد، فاصله دو جمله بالا یک پرش محاسبه می‌شد. همچنین کار روی زمان تعامل کاربر با سامانه نیز یکی از موضوعات حائز اهمیت دیگر است که باید در کارهای آتی مورد توجه بیشتری قرار گیرد.

منابع

Barrachina, S., Bender, O., Casacuberta, F., Civera, J., Cubel, E., Khadivi, S., Lagarda, A., Ney, H., Tomás, J., Vidal, E. and Vilar, J. M. (2007). Statistical Approaches to Computer-Assisted Translation, Computational Linguistics, Volume 35, pp. 3-28.

Bender, O., Hasan, S., Vilar, D., Zens, R. and Ney, H. (2005). Comparison of generation strategies for interactive machine translation, In Proceedings of the 10th Annual Conference of the European Association for Machine Translation (EAMT 05), pp. 33-40.

Brown, R.D., Nirenburg, S. (1990). Human-computer interaction for semantic disambiguation, In Processing of the International Conference on Computational Linguistics (COLING), PP. 42-47.

Brown, P. F., Della Pietra, S. A., Della Pietra, V. J. and Mercer, R. L. (1993). The mathematics of statistical machine translation: Parameter estimation, Computational Linguistics, pp. 263-311.

Civera, J., Vilar, J. M., Lagarda, A. L., Cubel, E., Barrachina S., Casacuberta, F., Vidal, E. (2006). Computer-Assisted Translation Tool based on Finite-State Technology. In: Proc. of EAMT 2006, pp. 33-40.

Cubel, E., González, J., Lagarda, A. L., Casacuberta, F., Juan, A. and Vidal, E. (2004). Adapting finite-state translation to the TransType2 project, Proceedings of the Joint Conference combining the 8th International Workshop of the European Association for Machine Translation.

Foster, G., Isabelle, P. and Plamondon, P. (1997). Target-Text Mediated Interactive Machine translation, in Kluwer Academic Publishers, pp. 175-194.

Foster, G. (2002). Text Prediction for Translators, Ph.D. thesis, Université de Montréal, Canada.

Kay, M. (1973). The MIND system, in Natural Language Processing, pp. 155-188.

Knight, K. (1999). Decoding complexity in word replacement translation models, Computational Linguistics, 25(4):607-615.

اما عملیاتی شدن این توانایی بالقوه، مستلزم دسترسی به فرهنگ‌های واژگان کامل است (برای محاسبه احتمال $p(e|f)$).

اگر در محاسبه احتمال مدل ترجمه IBM-1 برای یک کلمه نامزد e ، مدخل مرتبط با این کلمه و کلمات جمله مبدأ در فرهنگ واژگان موجود نباشد، مدل IBM-1 قادر به انتخاب این کلمه نخواهد بود. اگر وضعیت برای سایر کلمات ممکن دیگر نیز به همین صورت باشد، مدل IBM-1 هیچ پیشنهادی را نمی‌تواند به کاربر ارائه کند. مدل زبانی ۲-گرام تنها احتمال رخداد دو کلمه متوالی را بررسی می‌کند و محدودیت کمتری نسبت به مدل IBM-1 دارد و در هر شرایطی می‌تواند یک کلمه را به‌عنوان پیشنهاد به کاربر ارائه می‌دهد، هر چند که پیشنهادهای آن قطعیت کمتری نسبت به مدل IBM-1 دارد.

هدف ما از ترکیب این دو مدل، بالابردن توانایی ارائه پیشنهاد با قطعیت بیشتر بوده است. نتایج، موفقیت این روش را نشان می‌دهند. در هر دو روش جستجو، بهترین نتیجه با استفاده از روش ترکیب مدل‌های زبانی و IBM-1 به‌وجود آمده است؛ که البته از میان این دو نتیجه، روش جستجوی مبتنی بر مجموع وزن‌دار هزینه ترجمه و فاصله ویرایشی در نهایت نتیجه بهتری را ارائه داده است.

۵- نتیجه‌گیری و پیشنهاد زمینه‌های

پژوهشی آتی

هدف ما در این مقاله بهبود و توسعه یک سامانه مترجم بار تعاملی انگلیسی-فارسی بود. ما با ارائه یک روش جستجوی جدید، برای جستجوی پیشنهاد در گراف جستجو، موفق به کاهش بیشتر تلاش کاربر، برای تایپ ترجمه مورد نظرش، شدیم. همچنین با ترکیب روش‌های مختلف Back-off دوباره توانستیم به میزان بیشتری در کاهش تلاش کاربر مؤثر باشیم.

اما همچنان موضوعات دیگری در این پروژه وجود دارند که باید مورد تحقیق و بررسی قرار گرفته و بهبود داده شوند؛ یکی از این موضوعات، مسأله استفاده از معیار حداقل فاصله ویرایشی برای جستجوی پیشنهاد در گراف جستجو است. از آنجایی که این معیار جابه‌جایی‌های مجاز عناصر جمله را در نظر نمی‌گیرد، معیار زیاد مناسبی برای مقایسه فرضیه‌ها برای تطابق بیشتر با پیشنهاد نیست. برای مثال دو جمله "مولوی شاعر بزرگ ایران است" و "شاعر بزرگ ایران

Computational Linguistics (COLING), pages 329-334.

Zens, R., Och, F. J. and Ney, H. (2002). Phrase-Based Statistical Machine Translation, in Springer-Verlag Berlin Heidelberg, pp. 18-32.



زینب وکیل مدرک کارشناسی خود را در رشته مهندسی کامپیوتر از دانشگاه الزهرا (س) در سال ۱۳۸۷ و مدرک کارشناسی ارشد خود را در رشته هوش مصنوعی در سال ۱۳۹۱ از دانشگاه صنعتی امیرکبیر دریافت کرده است. زمینه‌های پژوهشی مورد علاقه ایشان پردازش زبان طبیعی و ترجمه ماشینی است. نشانی رایانامه ایشان عبارت است از:

z.vakil@aut.ac.ir



شهرام خدیوی مدارک کارشناسی و کارشناسی ارشد خود را در رشته مهندسی کامپیوتر از دانشگاه صنعتی امیرکبیر به ترتیب در سال‌های ۱۳۷۵ و ۱۳۷۸ دریافت کرده و همچنین مدرک دکترای خود را در سال ۱۳۸۷ از دانشگاه آخن RWTH در رشته علوم کامپیوتر دریافت کرده‌اند. ایشان در حال حاضر عضو هیئت علمی و استادیار دانشکده مهندسی کامپیوتر دانشگاه صنعتی امیرکبیر هستند. زمینه‌های پژوهشی مورد علاقه ایشان پردازش زبان طبیعی، ترجمه ماشینی آماری و یادگیری ماشین است. نشانی رایانامه ایشان عبارت است از:

khadivi@aut.ac.ir

Koehn, P., Hoang, H., Birch, A., Callison-Burch, C., Federico, M., Bertoldi, N., Cowan, B., Shen, W., Moran, C., Zens, R., Dyer, C. J., Bojar, O., Constantin, A. and Herbst, E. (2007). Moses: Open source toolkit for statistical machine translation, In ACL Demo and Poster Session, Available: <http://www.statmt.org/moses/>.

Koehn, P. (2009a). A Process Study of Computed Aided Translation, Kluwer Academic Publishers.

Koehn, P. (2009b). A web-based interactive computer aided translation tool, In Proceedings of the ACL Interactive Poster and Demonstration Sessions.

Koehn, P. (2010). Statistical Machine Translation, textbook, Cambridge University Press, page 203.

Levenshtein, V. I. (1965). Binary Codes Capable of Correcting Deletions, Insertions, and Reversals. Soviet Physics - Doklady, Vol. 10 No. 8 pp. 707-710.

Langlais, P., Foster, G., and Lapalme, G. (2000). TransType: a computer-aided translation typing system, In Proceedings of the NAACL/ANLP Workshop on Embedded Machine Translation Systems, pp. 46-52.

Langlais, P., Lapalme G. and Loranger, M. (2002). TRANSType: Development-Evaluation Cycles to Boost Translator's Productivity, in Kluwer Academic Publishers, pp. 77-98.

Maruyama, H., Watanabe, H. (1990). An interactive Japanese parser for machine translation, In Processing of the International Conference on Computational Linguistics (COLING), pp. 257-262.

Megerdounian, A., Khadivi, S. (2010). On evaluation of interactive-predictive machine translation: Practical problems and suggestions, presented at the Fifth International Symposium On Telecommunication (IST2010), Tehran, Iran.

Ney, H., Och, F., Vogel, S. (2000). Statistical Translation Of Spoken Dialogues In The Verbmobil System. In Workshop on Multi-Lingual Speech Communication , pages 69-74.

Ortiz-Martinez, D., García-Varea, I. and Casacuberta, F. (2010). Online Learning for Interactive Statistical Machine Translation, In The 2010 Annual Conference of the North American Chapter of the ACL, pp. 546-554.

Whitelock, P. J., McGee Wood, M., Chandler, B. J., Holden, N. and Horsfall, H. J. (1986). Strategies for interactive machine translation: the experience and implications of the UMIST Japanese project, In Proceedings of the International Conference on