



وزارت علوم، تحقیقات و فناوری

پژوهشگاه علوم و فناوری اطلاعات ایران

شناسایی طرح‌های مرتبط با تحقق

مترجم ماشینی در زبان فارسی

(گزارش نهایی)

سید مهدی سمائی

تیر ۱۳۹۰

فهرست مطالب

فصل اول: سابقه پژوهش ۲

- ۱-۱- مقدمه ۳
- ۲-۱- تعریف مسئله ۵
- ۳-۱- هدف ۶
- ۴-۱- پیشینه تحقیق ۷
- ۱-۴-۱- یک مرور کلی ۷
- ۲-۴-۱- بررسی روش‌ها و رهیافت‌ها ۱۱
- ۳-۴-۱- مترجم ماشینی ۲۲
- ۴-۴-۱- پژوهش‌های مرتبط با ترجمه ماشینی ۲۶
- ۵-۴-۱- جمع‌بندی ۴۶

فصل دوم: ترجمه ماشینی در زبان فارسی ۴۸

- ۱-۲- مقدمه ۴۹
- ۲-۲- زبان فارسی ۵۰
- ۱-۲-۲- خط فارسی ۵۰
- ۳-۲- ترجمه متن فارسی ۵۱
- ۱-۳-۲- استخراج واژه‌ها ۵۱
- ۲-۳-۲- واژگان ۵۲
- ۳-۳-۲- تحلیل صرفی و نحوی ۵۳
- جمع‌بندی ۱۰۴

واژه‌نامه انگلیسی به فارسی ۱۱۰

فهرست منابع ۱۱۴

- سایر منابع و مآخذ: ۱۱۷

فصل اول: سابقه پژوهش

زبان‌شناسی رایانه‌ای یکی از شاخه‌های زبان‌شناسی است که در آن فنون و مفاهیم علم رایانه را در تحقیقات حوزه‌های مختلف به کار می‌برند. ترجمه ماشینی یکی از زیرشاخه‌های زبان‌شناسی رایانه‌ای است که نام‌هایی چون ترجمه رایانه‌ای^۱ و ترجمه تعاملی^۲ نیز به آن اطلاق می‌شود. در این شاخه تلاش می‌شود چگونگی به کارگیری نرم‌افزار رایانه‌ای برای ترجمه متن یا گفتار از یک زبان طبیعی به زبان دیگر مورد مطالعه و بررسی قرار گیرد.

علیرغم همه محدودیت‌هایی که وجود دارد برنامه‌های ترجمه ماشینی در سراسر دنیا کاربرد دارند. از بزرگ‌ترین کاربران ترجمه ماشینی اتحادیه اروپا است. پروژه مولتو که توسط دانشگاه گوتنبرگ هدایت می‌شود، بالغ بر ۲,۳۷۵ میلیون یورو از اتحادیه اروپا بابت ایجاد یک ابزار ترجمه که بیشتر زبان‌های اروپایی EU را پوشش دهد، دریافت می‌کند.

گوگل مدعی است که موتور ترجمه جدیدش نتایج امیدبخشی به دست خواهد داد. این موتور با روش آماری در گوگل لنگوئج تولز برای ترجمه عربی - انگلیسی و

^۱ computer-aided translation

^۲ interactive translation

چینی - انگلیسی، از طرف موسسه ملی استاندارد و فن آوری نمره ۴۲۸۱/۰ را دریافت کرده است.

نیروهای نظامی ایالات متحده امریکا نیز سرمایه گذاری های چشم گیری روی مهندسی زبان های طبیعی انجام داده اند. دفتر فن آوری پردازش اطلاعات در "دارپا"^۱ از برنامه هایی چون تایدز و مترجم بابی لون استفاده می کند. نیروی هوایی امریکا قراردادی یک میلیون دلاری جهت ایجاد فناوری ترجمه امضا کرده است.

کاربردهای وسیع و متنوعی می توان برای ترجمه ماشینی تصور نمود. از شبکه های اجتماعی اینترنتی گرفته تا کاربرد آن در موبایل برای گویشوران زبان های مختلف که با یکدیگر تجارت می کنند و مسافرانی که زبان کشوری که به آن سفر کرده اند را نمی دانند.

¹ DARPA

۱-۲- تعریف مسئله

در طرح حاضر تلاش می‌شود یک طرح کلی برای ساخت مترجم ماشینی در زبان فارسی ارائه شود. لازمه ساختن نرم افزار ترجمه ماشینی در اختیار داشتن تصویری کلی از منطق ترجمه ماشینی و مراحل انجام آن است. علاوه بر این باید مشخص شود که در هر مرحله کدام بخش از زبان مقصد و مبدا مورد تحلیل دستوری قرار خواهد گرفت.

با توجه به بررسی مقدماتی پیشینه مطالعات و تحقیقات انجام شده در این زمینه، به نظر می‌رسد که طرح از پیش تعیین شده و واحدی برای ساختن مترجم ماشینی برای زبان فارسی وجود ندارد که در آن اجزای دستوری زبان فارسی مشخص شده باشد. با در دست داشتن دورنمایی از ترجمه ماشینی می‌توان ارزیابی درستی از پژوهش‌های موجود در این زمینه کرد.

این طرح در دو مرحله ارائه می‌شود. در مرحله اول ضمن یک مرور کلی بر سابقه پژوهش‌هایی که دیگر متخصصان این حوزه انجام داده‌اند به معرفی برخی نرم افزارهای موجود پرداخته خواهد شد.

در مرحله دوم طرح کلی ساخت نرم افزار مترجم فارسی ارائه خواهد شد.

۱-۳- هدف

ارائه طرح کلی ساخت یک مترجم ماشین در زبان فارسی

از سال ۱۹۵۴ به بعد بود که اصطلاح ماشین ترجمه به تدریج به کار رفت (دلونسی^۱، ۱۹۷۲)، ولی عبارت مترجم الکترونیکی و مترجم خودکار را هم گاهی به کار می‌بردند. اما تلاش برای ساختن مترجم ماشینی از دهه سی میلادی شروع شده بود. مترجم ماشینی را آن زمان وسیله ساده‌ای نظیر ماشین حساب می‌دانستند که به راحتی زبانی را به زبان دیگر ترجمه می‌کند. این ساده‌انگاری اندک‌اندک جای خود را به واقع بینی داد. یکی از محققان پس از دو دهه تلاش، ترجمه ماشینی را یکی از پیچیده‌ترین فعالیت‌هایی دانست که بشر تا آن زمان به آن دست زده است (دریفوس^۲: ۱۹۸۹: ۵۶). نشانه‌ها در ماشین‌های عددی و الفبایی، تک‌معنا و شفاف هستند و ابهام و شرح و تفسیر در آن‌ها جایی ندارد و اجتماع نشانه نیز ارزش معنایی واحد و مطلق دارند. در صورتی که در ترجمه زبان بشر، ابهام واژگانی^۳ و ابهام ساختاری^۴ از عمده مشکلات محققان هستند.

اولین مترجم ماشینی ساده و ابتدایی را یک محقق روسی در سال ۱۹۳۳ در مسکو ساخت و به ثبت رساند. در سال ۱۹۴۶ بوت^۵ و ویور^۶ به تحقیق درباره مترجم ماشینی پرداختند. ویور کار خود را بر اساس روش‌هایی که با آن، پیام‌های دشمن را در جنگ

^۱ Delavenay

^۲ Dreyfus

^۳ Lexical ambiguity

^۴ Lexical ambiguity

^۵ Boot

^۶ Weaver

جهانی دوم رمزشکنی می کردند قرار داده بود. بوت ماشینی ساخته بود که کلمات زبان را به زبان دیگر ترجمه می کرد. در کار بوت نه به نحو پرداخته می شد نه به ترتیب کلمات^۱. مترجم ماشینی بوت کلمات مهم متن یا کلیدواژه ها^۲ را به زبان دیگر ترجمه می کرد و فهم و تفسیر متن را به عهده کاربر وا می گذاشت.

ریچنز^۳ انگلیسی در سال ۱۹۴۸ فکر تجزیه دستوری خودکار پایانه کلمات^۴ را مطرح کرد که نوعی ترجمه کلمه به کلمه بود. ویور در سال ۱۹۴۹ ضمن تایید کار بوت و ریچنز شباهت های زیر ساختی زبان ها با یکدیگر را مطرح کرد؛ یعنی شباهت های معنایی و منطقی زبان ها که مربوط به ویژگی های مغز انسان هستند و در همه انسان ها مشترک اند. آنچه ویور در نهایت به کار بوت و ریچنز اضافه کرد پیشنهادهایی برای رفع ابهام معنایی بود.

در سال ۱۹۵۰ رایفلر^۵ نوعی ترجمه را که انسان با کمک ماشین انجام می داد مطرح کرد. به عقیده او باید ابتدا متن ترجمه می شد و سپس ماشین آن را ویرایش می کرد. یعنی به عبارتی، ترجمه ماشینی باید یک پیش ویرایش^۶ و یک پس ویرایش^۷ می- داشت و تعاملی^۸ می بود (بوت، ۱۹۶۷: ۵۳).

¹ Word Order

² Key words

³ Richens

⁴ Ending

⁵ Reifler

⁶ Pre-editing

⁷ Post-editing

⁸ Interactive

تحقیقات به صورتی گسترده ولی پراکنده ادامه داشت. اسوالد^۱ و فلچر^۲ در آلمان روی تجزیه نحوی جملات کار می کردند. بنیاد راکفلر در سال ۱۹۵۲ در نخستین همایش هجده نفری زبان شناسان و متخصصان رایانه، سرمایه قابل توجهی را برای ترجمه ماشینی به "موسسه فناوری ماساچوست" (ام آی تی) اهدا کرد. اعضای این همایش موافقت کردند که تحقیقات در دو مرحله انجام شوند: مرحله اول شامل تحقیق درباره بسامد کلمات و شیوه معادلیابی آنها و چگونگی استفاده از حافظه رایانه و دیگر جنبه های فنی فرهنگ لغت های ماشینی بود؛ مرحله دوم مربوط به شیوه تحلیل های نحوی می شد. از سال ۱۹۵۲ تا ۱۹۵۵ محققان آمریکایی درباره ظرفیت حافظه رایانه و شناسایی خودکار معنای کلمات و شیوه تحلیل پایانه های کلمات (پسوندها) تحقیق کردند. در سال ۱۹۵۴ دوسترت^۳ و گاروین^۴ در دانشگاه جورج تاون به موفقیت هایی در ترجمه روسی به انگلیسی دست یافتند. آن ها با ۲۵۰ کلمه و ۶ قاعده، الگویی برای ترجمه ماشینی عرضه کردند. در همان سال ام آی تی مجله "ماشین ترجمه" را منتشر کرد.

در سال ۱۹۵۶ ام آی تی اولین نشست بین المللی ترجمه ماشینی را با حضور متخصصان انگلیسی و کانادایی و آمریکایی و روسی برگزار کرد. در این زمان، آمریکا و انگلستان و شوروی محل اصلی تحقیقات ترجمه ماشینی بودند، ولی در ایتالیا و کشورهای اسکاندیناوی هم تحقیقاتی انجام می شد. محققان در این سال ها دیگر تلاش نمی کردند

^۱ Osvald

^۲ Fletcher

^۳ Dostert

^۴ Garvin

که ثابت کنند ترجمه ماشینی کاری شدنی است، بلکه در حال نظم دادن و متمرکز کردن پژوهش های پراکنده ای بودند که در نقاط مختلف جهان انجام می شد.

در سال ۱۹۵۸ در آمریکا دوازده گروه در کار پژوهش های مربوط به ترجمه ماشینی بودند: در دانشگاه هاروارد فرهنگ لغت ماشینی روسی به انگلیسی را می ساختند؛ چامسکی در ام آی تی درباره ساخت های نحوی تحقیق می کرد؛ در دانشگاه جرج تاون در باره نحو و معناشناسی زبان روسی پژوهش هایی انجام می شد و در دانشگاه میشیگان تلاش می کردند تا قاعده های نحو روسی را استخراج، و مشکل چند معنایی^۱ را حل کنند؛ در لس آنجلس و کالیفرنیا نیز محققان درباره مسائل روش شناختی ترجمه ماشینی تحقیق می کردند. در این بین به ترجمه ماشینی زبان روسی به انگلیسی بیشتر اهمیت داده می شد.

در انگلستان بیش تر درباره مسائل روش شناختی ترجمه ماشینی، تحقیقات دقیقی انجام شد. یکی از مهم ترین کارهای پژوهشگران آن کشور، درباره همکاری زبان شناسان و متخصصان رایانه بود. آنان تلاش بسیاری انجام دادند تا زبان شناسی، رشته ای از علوم دقیقه شود و در این راه از روش های ریاضی استفاده می کردند.

¹ Polysemy

از دهه پنجاه میلادی و با چامسکی، تحلیل های نحوی مستقل از معنا آغاز شده بود، ولی معنانشناسی را کتز و فودور (۱۹۶۳) و کتز و پستال (۱۹۶۴) وارد کار ترجمه ماشینی کردند. در مرحله بعد نیز کاربردشناسی وارد حوزه پردازش زبان شد.

انسان دانشی از عالم خارج دارد که آن را در تعبیر متن به کار می بندد. از این رو نظریه های گوناگونی درباره شکل تجربه انسان از عالم خارج عرضه شد. یعنی متخصصان سعی کردند دانش عالم خارج را در حافظه رایانه بگنجانند و عاقبت به این نتیجه رسیدند که بهتر است دانش گسترده عالم خارج را محدود کنند. به این سبب به جای ذخیره کل دانش در حافظه رایانه، موضوع عالم فرعی^۱ و نمایشنامه^۲ و فیلمنامه^۳ مطرح شد که همگی، شکل های مختلف ذخیره دانش عالم خارج در رایانه هستند.

۱-۴-۲- بررسی روش ها و رویافت ها

همان طور که گفته شد در سطح ابتدایی ابتدا به جایگزین سازی^۴ واژه های زبان مبدا با واژه های زبان مقصد پرداخته می شود. اما این جایگزینی به تنهایی ترجمه مناسبی به دست نمی دهد چرا که تشخیص گروه ها و و اجزای آن ها ضروری است. یکی از راه حل هایی متخصصان به کار گرفتند استفاده از پیکره^۵ و تکنیک های آماری بود.

¹ Subworld

² Scripts

³ Scenario

⁴ substitution

⁵ corpus

با محدود کردن قلمرو، مثل ساخت یک مترجم در یک حرفه مشخص – مثل نرم افزار گزارش وضع آب و هوا – می توان امکان جایگزین سازی واژه ها را محدود کرد. این شیوه بویژه در مواردی که یک زبان رسمی یا منضبط^۱ به کار می رود، مفید واقع می شود. به همین دلیل است که ترجمه ماشینی اسناد دولتی و حقوقی در مقایسه با گفتگوها یا متونی که کم تر استاندارد هستند، نتایج بهتری در بردارد.

هم چنین با دخالت انسان در ترجمه می توان انتظار نتایج بهتری داشت. برای مثال چنان چه کاربر بتواند دقیقاً تعیین کند کدام واژه ها اسم هستند، نرم افزار قادر به ترجمه دقیق تری خواهد بود. با استفاده از چنین تکنیک هایی، ترجمه ماشینی ثابت کرده است که می تواند به عنوان یک ابزار مفید جهت کمک به ترجمه توسط انسان به کار گرفته شود و حتی می تواند خروجی های قابل استفاده ای تولید کند (برای مثال گزارش وضع آب و هوا).

می توان گفت که فرآیند ترجمه شامل دو مرحله رمزگشایی معنای متن از زبان مبدا و به رمز درآوردن همان معنا در زبان مقصد است. در پس این فرآیند به ظاهر ساده یک عملیات شناختی پیچیده نهفته است. مترجم به باید تمام جنبه های زبان را مورد تحلیل و تفسیر قرار دهد. این فرآیند نیاز به برخورداری از دانشی عمیق در مورد جنبه های

¹ formulaic

مختلف هر دو زبان (زبان مبدا و مقصد)، مثل صرف و نحو، معنی‌شناسی، اصطلاحات و... و حتی مسائل فرهنگی گویشوران آن زبان‌ها دارد.

باتوجه به آن‌چه که گفته شد یک چالش مهم در ترجمه ماشینی این است که چگونه می‌توان برنامه‌ای طراحی کرد که رایانه متن را مثل یک انسان درک کند، و مثل یک انسان متنی جدید در زبان مقصد تولید کند. رهیافت‌های مختلفی برای حل این مسئله وجود دارد.

ترجمه ماشینی می‌تواند از روشی مبتنی بر قواعد زبان‌شناختی بهره بگیرد. بدین معنی که واژه‌ها به طریقی زبان‌شناختی ترجمه شوند- بهترین ترجمه در زبان مقصد جایگزین واژه‌ها در زبان مبدا شوند.

به عقیده برخی از متخصصان "درک متن" مسئله‌ای مقدم بر ترجمه ماشینی است که موفقیت ترجمه ماشینی در گرو حل آن است.

براساس روش‌های قاعده بنیان^۱، یک میانجی^۲ لازم است که متن ابتدا با آن بازنمایی شده و سپس به زبان مبدا ترجمه می‌شود. بر اساس ماهیت این میانجی، رهیافت "ترجمه ماشینی میان‌زبانی"^۳ یا "ترجمه ماشینی انتقال بنیان"^۴ نامیده می‌شود. این روش‌ها نیازمند واژگان وسیع، اطلاعات صرفی، نحوی و معنایی و نیز تعداد زیادی قاعده است.

^۱ rule-based

^۲ intermediary

^۳ interlingua machine translation

^۴ transfer-based machine translation

معمولا گویشور بومی یک زبان، در صورتی که داده کافی وجود داشته باشد، از ترجمه یک مترجم ماشینی می تواند به طور تقریبی به منظور گویشور دیگر پی ببرد. در این روش مسئله فراهم کردن داده کافی و مناسب است. به طور مثال برای روش های آماری به حجم عظیمی از پیکره های چندزبانه نیاز است.

برای ترجمه زبان هایی که با یکدیگر ارتباط نزدیکی دارند می توان از روشی به نام "ترجمه ماشینی سطحی"^۱ استفاده نمود.

روش های قاعده بنیان شامل سه روش انتقال بنیان، میان زبانی و فرهنگ بنیان^۲ می - شود.

ترجمه ماشینی میان زبانی : یکی از انواع ترجمه ماشینی قاعده بنیان، ترجمه ماشینی میان زبانی است. در این رهیافت زبان مبدا به یک زبان میانجی تبدیل می شود. سپس زبان مقصد از این زبان میانجی تولید می شود.

ترجمه ماشینی فرهنگ لغت بنیان: این ترجمه بر اساس مدخل های فرهنگ لغت است. در این روش واژه ها مطابق با مدخل فرهنگ لغت ترجمه می شوند. در این روش از پیکره های متنی دوزبانه مانند کنیدین هنسرد^۳ گزارش انگلیسی - فرانسوی پارلمان کانادا، و یوروپال^۴ گزارش پارلمانی اروپا استفاده می شود. در شرایطی که چنین پیکره هایی در اختیار باشد ترجمه متون مشابه نتایج قابل توجهی دربر خواهد داشت، اما به هر حال چنین

^۱ shallow-transfer machine translation

^۲ dictionary-based

^۳ Canadian Hansard

^۴ EUROPARL

منابعی بسیار محدود است. اولین مترجم ماشینی آماری کاندید^۱ از آی بی ام^۲ بود. گوگل برای سال‌های متمادی از سیستم^۳ استفاده کرد اما در اکتبر ۲۰۰۷ به یک روش ترجمه آماری روی آورد. گوگل با افزودن حدود ۲۰۰ بلیون واژه دقت ترجمه خود را بهبود بخشید.

ترجمه ماشینی مثال بنیان^۴: این روش توسط ماکوتوناگو در ۱۹۸۴ بنیان نهاده شدⁱⁱⁱ. در این روش از یک پیکره دوزبانه به عنوان منبع دانش اصلی، استفاده می‌شود. در این روش ترجمه بر اساس قیاس^۵ انجام می‌شود.

ترجمه ماشینی پیوندی^۶: این روش ترکیبی است از روش‌های قاعده بنیان و آماریⁱⁱⁱ. برخی شرکت‌های ترجمه ماشینی مانند ایژیا آن‌لاین^۷، لینگواسیس^۸، سیستم^۹ یوپی وی مدعی استفاده از چنین روش‌هایی هستند.

این روش با رویکردهای مختلفی می‌تواند به کار رود:

پس‌ویرایش آماری قواعد^{۱۰}: ترجمه توسط یک موتور قاعده بنیان انجام می‌شود. سپس از آمار برای اصلاح خروجی این موتور استفاده می‌شود.

¹ CANDIDE

² IBM

³ SYSTRAN

⁴ Example-based machine translation

⁵ analogy

⁶ Hybrid machine translation

⁷ Asia Online

⁸ LinguaSys

⁹ UPV

¹⁰ Rules post-processed by statistics

ویرایش آمار با قواعد^۱: به منظور بهبود عملکرد موتور آماری، داده‌ها با استفاده از قواعد پیش‌پردازش می‌شوند. هم‌چنین خروجی موتور آماری نیز توسط قواعد پس‌پردازش می‌شود. این رهیافت از انعطاف و قدرت و کنترل بالایی برخوردار است.

۱-۴-۲-۱- برخی مسائل و مشکلات

امروزه با وجود پیشرفت‌هایی که در ترجمه ماشینی حاصل شده، هنوز مشکلات رایجی مثل هم‌معنایی^۲، چند معنایی، ابهام ساختاری، ابهام واژگانی، و حتی محدود کردن عالم خارج و مشخص کردن بافت^۳ به جا مانده است. در ابهام‌زدایی از معنای واژه^۴ (WSD) به دنبال یافتن مناسب‌ترین معنی برای واژه‌هایی که به چند معنای مختلف دلالت می‌کنند، هستیم. این مسئله اولین بار در سال‌های ۱۹۵۰ توسط یهوشوا بارهیل^۵ مطرح شد^۶. او اعلام کرد که بدون یک دایره‌المعارف جهانی^۷ یک مترجم ماشینی هرگز قادر به تمیز دادن بین دو معنای مختلف یک واژه نخواهد بود^۸. امروزه رهیافت‌های گوناگونی برای حل این مشکل به کار بسته می‌شوند که به طور کلی می‌توان آن‌ها را به دو دسته سطحی^۷ و عمیق^۸ تقسیم کرد. به طور کلی می‌توان گفت در دسته اول با استفاده از روش‌های

^۱ Statistics guided by rules

^۲ Homonymy

^۳ Context

^۴ Word-sense disambiguation

^۵ Yehoshua Bar-Hillel

^۶ universal encyclopedia

^۷ shallow

^۸ deep

آماری، از طریق واژه‌های همسایه واژه مبهم عمل می‌شود و دانش متن در نظر گرفته نمی‌شود. و این درست برعکس روش عمیق است که به میزان گسترده‌ای دانش متن در نظر گرفته می‌شود.

از دیگر مشکلات که به ویژه در زبانی چون فارسی مشکل آفرین است، آزاد بودن ترتیب کلمات (آرایش واژه‌ها) است. به این معنی که فارسی زبانان مخیرند سازه‌های مستقل را جابه‌جا کنند بدون آنکه به ویژگی دستوری بودن کلمات آسیبی وارد شود. ضمناً رسم‌الخط فارسی به دلایل مختلف در کار ترجمه ماشینی ممکن است اخلال ایجاد کند.

همانطور که گفته شد برای ساختن مترجم ماشینی باید عالم خارج را محدود کرد. یکی از مشکلاتی که مشخص نبودن بافت بوجود می‌آورد ابهام واژگانی است. گاهی می‌توان با توجه به دانش زبانی ابهام را برطرف کرد. مانند جمله زیر:

۱- Jane will coach our team.

در جمله شماره ۱ لفظ coach فعل است و در این بافت فقط فعل می‌نشیند و مشکلی پیش نمی‌آید. اما اگر دو اسم هماوا^۱ با دو معنی متفاوت در دو بافت متفاوت قرار گیرند ابهامی پیش می‌آید که نمی‌توان آن را با دانش زبانی^۱ برطرف کرد:

¹ homophone

Two teams were trained by the same coach. –۲

The coach was heading for the station. –۳

در جمله شماره ۳ لفظ coach به معنی وسیله نقلیه است که جز با آگاهی از بافت نمی توان به معنی آن در جمله مذکور پی برد. (روزتا، ۱۹۹۴: ۳).

ابهام ساختاری نوعی دیگر از ابهام است که به دلیل مشخص نبودن بافت پیش می آید. در جمله شماره ۴ بدون کمک بافت نمی توان فهمید که آیا دختر دوربین دارد یا مرد:

The girl followed the man with the camera. –۴

یکی از کاربردهای دستورسنج برای امتحان صحت دستوری یک زبان واحد است. یک نمونه از این دستورسنج ها را دانشکده زبان های دانشگاه هایدلبرگ آلمان در چارچوب طرحی به نام ویرایشگر خود کار پلین^۲ ساخته است. این طرح از آوریل ۱۹۸۹ تا مارس ۱۹۹۴ طول کشیده و مجری آن پیترهلویگ^۳ بوده است (لیزا، ۱۹۹۷).

نرم افزار دیگری به نام ورپرفکت^۴ وجود دارد که ترکیب های دستوری و غیردستوری زبان انگلیسی را از هم تمیز می دهد. نرم افزار مذکور الگوی این رساله است.

شیوه کار ورد پرفکت صوری و فارغ از بافت است. این نرم افزار می تواند ایرادهای نحوی و املائی را تشخیص دهد. در صورت مشاهده اشکال در ساخت دستوری جمله پیشنهاد

¹ linguistics knowledge

² Maschinelle Syntaxprüfung und Korrektur PLAIN – C.

³ Peter Hellwig

⁴ word perfect

شده، نوع اشکال را مشخص می کند و اگر املاي کلمات نادرست باشد تذکر می دهد و کلمه های هم خانواده با کلمه نادرست را جدولی فهرست می کند تا یکی از آن میان انتخاب شود. در این جا برای نشان دادن شیوه عملکرد نرم افزار چند نمونه از داده هایی که به برنامه ارائه شده جواب های مربوط به آن ذکر می شود.

پاسخ نرم افزار	داده‌ها
Grammatical	John eats the apple -۱
Grammatical	The apple is going to school -۲
A compound subject requires a plural verb.	John and Mary eats the apple -۳
"The" doesn't usually come before John.	The John eats the apple -۴
Grammatical	John eats -۵
The noun phrase appears incomplete.	John eats the -۶
This doesn't seem to be a complete sentence. An article or other modifiers usually precedes the word "book".	In book -۷
This noun phrase appears incomplete. Check for word formation and punctuation errors for missing word.	I:my -۸
Grammatical	John gave -۹
Grammatical	Apple -۱۰
The subject pronoun "who" should not be used in the object position.	The man who you saw eats an apple -۱۱
The singular subject "man" takes a singular verb not the plural verb "eat".	The man who eats the apple eat the dog -۱۲
Spelling error : eaten, eater, eat, ...	John eats the apple -۱۳
Grammatical	The man that -۱۴

از داده‌های نرم‌افزار به نتایجی می‌توان دست یافت:

الف. از داده شماره ۲ درمی‌یابیم که نرم‌افزار ساخت نحوی ساخت نحوی را ملاک قرار می‌دهد نه معنی را.

ب. از داده شماره ۳ و ۱۲ می‌توان اولاً به قابلیت نرم‌افزار در تشخیص نهاد مفرد از نهاد جمع پی برد و ثانیاً به این نکته رسید که داده‌ها فقط با نظریه مارکف^۱ به صورت خطی^۲ تحلیل نمی‌شوند و ساختهای نظیر جمله‌های پایه و پیرو از هم تمیز داده می‌شوند.

ج. از داده شماره ۶ و ۷ و ۸ می‌توان دریافت که قاعده گروه‌ها و نرم‌افزار وجود دارد.

د. داده شماره ۱۱ ثابت می‌کند که نرم‌افزار نقش ساختاری فاعل و مفعول را می‌شناسد.

ه. از داده شماره ۸ می‌توان به این نکته نیز پی برد که قاعده‌های علائم سجاوندی^۳ در نرم‌افزار تعبیه شده است.

و. داده شماره ۵ و ۹ ثابت می‌کند که قواعد زیرمقوله‌ای تعیین نشده است.

ز. از داده شماره ۱۰ و ۱۴ می‌توان فهمید که نرم‌افزار داده‌ها را با قاعده‌های جمله‌سازی انطباق نمی‌دهد.

لازم به ذکر است که در این رساله قاعده‌های جمله‌سازی و قاعده‌های زیرمقوله‌ای ارائه خواهد شد.

^۱ Markov

^۲ linear

^۳ punctuation

یکی از مترجم‌های ماشینی که در آن موضوع ترجمه محدود شده و در واقع بافت مشخص شده است برنامه تاوم متئو^۱ است که در هواشناسی کانادا از آن استفاده می‌کنند. این برنامه که ساخت دستوری آن نیز محدود است می‌تواند اخبار هواشناسی را به صورت کاملاً خودکار از زبان انگلیسی به زبان فرانسوی ترجمه می‌کند (روزتا: ۷).

نمونه دیگر وی. ال. اس^۲ نام دارد که در سال ۱۹۹۲ ساخته شده و با پیش‌ویرایش عمل می‌کند. یعنی شخصی که محدودیت‌های مترجم خودکار را می‌شناسد متن را قبل از وارد کردن در دستگاه تنظیم می‌کند و در واقع نوعی زبان مهارشده^۳ بر اساس بافت می‌سازد (روزتا: ۸).

از دیگر مترجم‌های خودکار یکی سیستم سیسترن است که در سال ۱۹۹۲ و دیگری متال^۴ است که در سال ۱۹۸۴ به کار گرفته شده است. این دو مترجم فقط قادر به ترجمه متون فنی هستند و هر دو به صورت پس‌ویرایش عمل می‌کنند. یعنی پس از ترجمه با دست اصلاح می‌شوند (روزتا: ۸).

روزتا دیگر مترجم خودکاری است که از سال ۱۹۸۵ تا ۱۹۹۲ در شهر انده‌وون^۵ و در پژوهشکده فیلیپس^۶ ساخته شده است.

^۱ TAUM METO

^۲ V.L.S

^۳ controlled language

^۴ METAL

^۵ Eindhoven

^۶ Philips Research Laboratories

نرم افزاری به نام پاور ترانسلیتور^۱ نیز در ایالات متحده ساخته شده که در ایران نیز در دسترس هست و می تواند متون انگلیسی و فرانسوی و آلمانی و اسپانیایی را به یکدیگر تبدیل کند. این مترجم خود کار محدودیت بافت ندارد و از این رو ابهام واژگانی و ساختاری باعث اخلاص در عملکرد آن می شوند.

همه این مترجم ها دارای واژگانی هستند که در آن ها ویژگی های صرفی و نحوی و معنایی هر واژه ثبت شده است. واژه ها بسته به مقولاتی که دارند دارای حوزه قواعد نیز هستند.

امروزه نرم افزارهای ترجمه زیادی وجود دارد که برخی از آن ها به صورت آن-لاین هم کار می کنند. برخی از آن ها عبارت اند از:

- آنوسارکا^۲: یک مترجم رایگان و متن باز^۳ انگلیسی به هندی است که براساس دستور پانینی^۴ کار می کند.
- آپ تک^۵: که سیستم ترجمه پیوندی را در ۲۰۰۹ عرضه کرد.
- عربیک ماشین ترانسلیشن^۶: یک مترجم چندزبانه است.
- آسیا آن لاین^۷: یک موتور ترجمه است که ابزارهایی برای ویرایش متن هم دارد.
- گوگل ترانسلیت^۸: یک مترجم آلاین از سایت گوگل است.

^۱ Power Translator

^۲ Anusaaraka

^۳ open source

^۴ Panini

^۵ AppTek

^۶ Arabic machine translation

^۷ Asia Online

^۸ Google Translate

- گوگل ترانسلیتور تول کیت^۱: وبسایتی برای ارائه خدمات جهت ویرایش متن‌هایی است که توسط گوگل ترجمه شده‌اند.
- مترجم ماشینی هندی به پنجابی^۲
- ایدیومکس^۳
- لینگواسیس^۴: یک مترجم با استفاده از روش پیوندی است که می‌تواند از هر زبانی به زبان دیگر برود.
- نروترن^۵: نرم‌افزاری است که قادر به ترجمه کتاب، صفحات وب، ایمیل، فکس، یادداشت، کتاب راهنما، گزارش، نامه و ... به زبان‌های مختلف است.
- پرپیت^۶
- شی شی ترا^۷: یک مترجم اسپانیایی - کاتالونیایی
- تاگل تکست^۸: مترجم انگلیسی - اندونزیایی با استفاده از روش انتقال بنیان
- ورلد لینگو^۹: هم آمار و هم قاعده بنیان است. در مایکروسافت ویندوز و مک آفیس^{۱۰} از آن استفاده می‌شود.
- دیلماج

¹ Google Translator Toolkit

² Hindi to Punjabi Machine Translation System

³ IdiomaX

⁴ LinguaSys

⁵ NeuroTran

⁶ Prompt

⁷ SiShiTra

⁸ Toggletext

⁹ Worldlingo

¹⁰ Mac Office

هیچ‌یک از مترجم‌های موجود ترجمه کامل و بدون نقص و کاملاً خودکار از

متن‌های محدود نشده^۱ ارائه نمی‌دهد^{viii vii vi}. اما به هر حال کیفیت ترجمه وقتی قلمرو^۲

محدود می‌شود به شکل قابل توجهی بهبود می‌یابد^{ix}.

^۱ unrestricted

^۲ domain

طیبی (۱۳۷۴) در رساله کارشناسی ارشد خود که عنوان آن "کاربرد دستور واژگانی نقشمند در ترجمه ماشینی" است، پژوهشی نظری برای ترجمه برخی تلکس‌های موجود در سازمان هواپیمایی کشوری کرده است. تحقیق او بر اساس نظریه برزنن^۱ یعنی دستور واژگانی نقشمند^۲ بوده است.

ثانی (۱۳۷۸) در مقاله‌ای با نام "ترجمه نحوی انگلیسی به فارسی با روش تجزیه چارت" به یک معرفی کلی از روش تجزیه چارت می‌پردازد و سپس مترجمی که آن را

LANT نامیده تشریح میکند. ثانی درباره روش تجزیه چارت به نقل از آلن^۳ (۱۹۹۵) و کرولی^۴ (۱۹۹۱) می‌نویسد که این روش یکی از سریع‌ترین روشهای تجزیه نحوی موجود (است که همه تجزیه‌های ممکن یک جمله را به طور همزمان به دست می‌آورد و عقبگرد نمی‌کند. در این روش برای جلوگیری از تحلیل مجدد سازه‌های مشابه، از یک ساختمان داده‌ای خاص به نام چارت استفاده میشود. تمام سازه‌های مجاز که در بین دو نقطه از یک جمله پیدا میشوند، در چارت نگهداری میشوند تا در صورت لزوم دوباره مورد استفاده قرار بگیرند.

^۱ Bresnan

^۲ lexical Functional Grammar

^۳ J. Allen

^۴ G. Krulee

در این مقاله عنوان میشود که در LANT ضمن به کارگیری روش چارت، سعی بر این بوده است که دانش ترجمه به گونه‌ای نمایش داده شود که بتوان با معرفی دانش ترجمه زبان‌های مورد نظر، ترجمه از هر زبان دلخواه مبدائی را به یک زبان دلخواه مقصد انجام داد. در این سامانه، دانش ترجمه با زبانی به نام زبان بازنمایی دستور 1 بیان می‌شود. به نظر نویسندگان این مقاله با استفاده از این زبان می‌توان قواعد دستوری یک زبان را در واحدهایی به نام دستور، واژگان و ساختواره دقیقاً توصیف نمود. LANT با استفاده از قواعد نحوی یک زبان طبیعی که با استفاده از زبان بازنمایی دستور بیان شده‌اند، جمله‌های آن زبان را از نظرنحوی تجزیه و تحلیل می‌کند. و برای ترجمه باید علاوه بر قواعد نحوی زبان مبدأ، یک طرح ترجمه نیز به کمک زبان بازنمایی دستور تعریف نمود. قواعد ترجمه‌ای این زبان مشابه قواعد گشتاری دستورهای گشتاری است. مترجم قادر است جمله‌های ساده LANT انگلیسی را به فارسی ترجمه کند. نیز این مترجم قادر به ترجمه برخی از الگوهای پیچیده نظیر بند موصولی هم هست.

شمس فرد (۱۳۷۴) نیز در رساله کارشناسی ارشد مهندسی نرم افزار طرحی برای درک متن^۱ ارائه کرده است. سیستم دنا قادر به درک متون (محدود شده) فارسی و تبدیل آن‌ها به بازنمایی داخلی با شبکه وابستگیهای مفهومی بر اساس نظریه شنک می‌باشد. دنا همچنین قادر به استنتاج نکات ضمنی متن و افزودن این دانش به بازنمایی داخلی است. به

¹ text understanding

این ترتیب می تواند با درک پرسش کاربر پاسخ وی را بر اساس آن چه صراحتاً در متن ذکر شده و یا از متن و دانش پیش زمینه سیستم قابل استنتاج است، ارائه نماید. (شمس فرد: ۱۳۸۵)

طرح پیشنهادی شمس فرد یک "واژه نامه" دارد که به گفته او "انتظارات" این واژه نامه با "حفره های خالی قاب افعال" پر شده اند. وی با استفاده از "نبشته ها"^۱ سعی کرده است دانش عالم خارج را برای درک متن به کار بندد. او از نظریه وابستگی مفهومی^۲ شنک^۳ برای ارتباط معنایی متن استفاده کرده است.

چنان که مشاهده می شود در رساله مذکور از اصطلاحات رایج زبان شناسی استفاده نشده است. واژه نامه او همان واژگان و حفره های خالی قاب افعال همان چارچوب زیرمقوله ای^۴ و انتظارات همان قواعد زیرمقوله ای^۵ هستند.

دو رساله از محمودی (۱۹۹۰ و ۱۹۹۴) در خارج از کشور ارائه شده است. یکی از این دو رساله در مقطع کارشناسی ارشد و دیگری در مقطع دکتری تهیه شده است. از آن جایی که محتویات رساله کارشناسی ارشد او در رساله دکتری اش وجود دارد فقط رساله دکتری او بررسی خواهد شد. موضوع تحقیق محمودی درباره پردازش گروه اسمی زبان فارسی و نوشتن قاعده های آن بر اساس برنامه پرولوگ^۶ است. او ابتدا به تعیین اجزای

^۱ به جای نبشته ها لفظ نمایشنامه در این رساله به کار رفته است.

^۲ conceptual dependency

^۳ Schank

^۴ subcategorization frame

^۵ subcategorization rules

^۶ Prolog

کلام پرداخته و قائل شدن به طبقای جدید را مطرح کرده است. این طبقات عمدتاً باید بر اساس توزیع^۱ تعیین شوند. وی سپس به پانزده طبقه واژگانی - نحوی^۲ قائل می شود که هر طبقه نیز به طبقاتی فرعی تقسیم می شود. محمودی در نهایت به ساختار گروه اسمی می رسد و الگوریتم های گروه اسمی را ارائه می کند.

شهابی (۱۳۷۶) پس از تحلیل لغوی، ساختوازی و نحوی، جمله های فارسی در محدوده معنایی خاص (درخواست اشتغال به کار افراد برای یک سازمان) را با شبکه معنایی بازنمایی می کند. (شمس فرد: ۱۳۸۵)

داوودآبادی و پالهنک (۱۳۸۴) سیستم دیگری در زمینه پردازش معنایی جملات فارسی ارائه کرده اند. این نرم افزار قادر به درک تعداد محدودی از جملات فارسی است. (شمس فرد: ۱۳۸۵)

شمس فرد (۱۳۸۱) یک سیستم یادگیر هستان شناسی از روی متون ساده زبان فارسی، به نام هستی ارائه کرده است. این سیستم پس از تجزیه نحوی جملات، نقش های موضوعی آن ها را استخراج کرده و بر اساس آن ها دانش مفهومی موجد در متن به صورت عناصر هستان شناسی به زبان "کیف"^۳ بازنمایی و به یک هسته اولیه هستان شناسی افزوده می - شود. (شمس فرد: ۱۳۸۵)

^۱ distribution

^۲ Lexico - syntactic

^۳ kif

رساله دکتری سیدمهدی سمائی (۱۳۷۷) با عنوان "واژگان در دستورسنج" انگاره نظری برنامه‌ای رایانه‌ای است که با آن می‌توان ترکیب‌های دستوری و غیردستوری زبان فارسی را از هم تمیز داد. به این منظور نگارنده قاعده‌های صوری ساخت جمله‌های خبری ساده مثبت را ارائه کرده است. این تلاش بر اساس فرضیه استقلال نحو چامسکی و الهام گرفتن از نرم‌افزار وردپرفکت انجام شده است. فرض رساله بر این بوده است که این دستگاه باید بتواند ترکیب‌های دستوری و غیردستوری را تمیز دهد و به محدودیت‌های باهم‌آیی ترکیب‌ها توجه نکند. در این کار محدودیت‌ها و قیودی ایجاد می‌شود که برخی از آن‌ها به شرح زیر است: (همان: ۱۳۸)

ناپایداربودن برخی مقوله‌های واژگانی یکی از مشکلات است. این که برخی واژه‌ها به یک مقوله واژگانی ثابت تعلق ندارند. برای مثال "خوب" متعلق به سه طبقه و "عاقلا نه" متعلق به دو طبقه واژگانی مختلف است. این مساله در تحلیل جمله‌هایی از این دست مشکل ایجاد می‌کند:

۱- سیب‌های تمیز و زیبا خوب شسته شدند.

در قواعد آوایی قید شده که اگر پس از صفت‌های مختوم به "ا" مثل "زیبا" صفت دیگری واقع شود باید یک همخوان میانجی "ی" پس از آن قرار گیرد. در جمله مذکور لفظ خوب پس از زیبا قرار گرفته ولی بین آن‌ها همخوان میانجی نیامده است. ماشین مفروض ممکن است این ترکیب را نادرست بداند و از تشخیص "خوب" در نقش قید

عاجز باشد. فرض کنیم در قواعد حوزه صفت آمده باشد که توالی بیش از دو صفت در گروه اسمی مجاز نیست. طبق این قاعده وقتی دستگاه به "زیبا" برسد این ایراد را مطرح می‌کند که صفت نمی‌تواند خارج از گروه اسمی و به صورت مستقل به کار رود و ترکیب غیردستوری است.

یکی از راه‌های مقابله با عناصر مترادف آوردن هم‌نویسه برای این گونه عناصر است. دستگاه با دیدن شماره هم‌نویسه در جست‌وجوس عناصر مترادف برمی‌آید. مثلاً دو "خوب" در حوزه واژگان هست که یکی صفت و دیگری قید است. با جست‌وجوی رفتار نحوی "خوب" در حوزه قید و انطباق آن با جمله سابق‌الذکر، دستوری بودن کاربرد "خوب" در آن بافت توجیه می‌شود.

یکی دیگر از مشکلات مربوط به تشخیص مرز گروه‌ها از یکدیگر است و در سازه‌یابی^۱ جمله زیر چنین مسأله‌ای می‌تواند رخ دهد:

۲- پسرهای خوب دانشجویان شهر سیب را خوردند.

مرز گروه اسمی در این جمله پس از "خوب" و "دانشجویان" و "شهر" و "را" می‌تواند باشد:

[[پسرهای خوب]_۱ [دانشجویان]_۲ [شهر]_۳ [سیب]_۴ را] خوردند

¹ parsing

چنانچه مرز گروه اسمی شماره ۱ و ۲ و ۳ باشد شناسه باید به صورت جمع بیاید ولی اگر شماره ۴ (پس از را) باشد شناسه می تواند جمع نیا مفرد باشد و سازه یابی هایی از این دست در نهایت می تواند تعاملی^۱ - یعنی با دخالت انسان - باشد.

نگارنده این رساله در نهایت اظهار داشته است که تلاش این رساله نشان دادن انگاره ای نظری برای نشان دادن ساز و کار دستورسنج بوده است و نیز نشان دادن مسائلی که در ارتباط با زبان در این کار پیش می آید.

"مترجم ماشینی فارسی - انگلیسی : یک نگاه کلی به پروژه شیراز"^۲ عنوان مقاله ای است که در آن پروژه ترجمه ماشینی شیراز معرفی شده است. این مقاله در ابتدا یک نگاه کلی به ویژگی های زبان شناختی زبان فارسی دارد. سپس به شرح مدل شیراز می - پردازد.

نویسندگان این مقاله پروژه شیراز را به عنوان یک پیش نمونه برای ترجمه فارسی به انگلیسی معرفی می کنند. این پروژه از اکتبر ۱۹۹۷ تا آگست ۱۹۹۹ به طول انجامیده است. این مترجم مجهز به یک فرهنگ لغت الکترونیکی دو زبانه فارسی - انگلیسی شامل حدود ۵۰۰۰۰ لفظ و یک تحلیل گرسرفی و یک تجزیه گرنحوی است. هدف از ساختن

¹ interactive

² Persian-English Machin Translation: An Overview of the Shiraz Project

این مترجم بیش تر ترجمه اخبار بوده است. این سیستم قادر به واحدسازی^۱، تشخیص کلمات مرکب و فعل سبک^۲ است. تجزیه گر نحوی قادر به تحلیل گروه اسمی، گروه حرف اضافه‌ای و بند موصولی است.

این مقاله ابتدا اطلاعاتی در مورد زبان فارسی شامل خانواده‌ای که به آن تعلق دارد، نگارش و رسم الخط فارسی ارائه می‌دهد. و سپس به طور پراکنده به مسائلی چون ابهام در صرف (به هم‌نویسه‌ها اشاره می‌کند، که البته با جمله "تکواژهای مختلفی که صورت‌های یکسانی دارند" به هم‌نویسه‌ها اشاره می‌کند)، "ساختار فعل‌های سبک"^۳ (که در واقع منظور فعل مرکب است) و مسائلی در نحو می‌پردازد. این مسائل عبارت‌اند از حالت و حالت‌نما، این که در زبان فارسی به استثنای مفعول‌نمای "را" حالت‌نمای مشخصی وجود ندارد و از این رو تعیین مرز واژه دشوار است. هم‌چنین این که هیچ نشان-گر خاصی که ارتباط بین سازه‌های یک گروه اسمی را مشخص کند وجود ندارد. تنها عنصری که ممکن است چنین نقشی را بازی کند تکواژ اضافه /e/ است که صورت نوشتاری ندارد (مگر این که بعد از مصوت‌های /â/ و /u/ بیاید و به صورت واج میانجی

¹ tokenization

² Light verb

³ Light verb construction

/y/ به نگارش دربیاید). ابهام در مرز^۱ گروه‌ها نیز مسئله دیگری است که در این قسمت به آن اشاره شده است. تجزیه گرنحوی باید مرز گروه‌ها (مثلاً گروه اسمی) و نیز نهاد و گزاره را مشخص کند. سپس به چند مورد که می‌تواند به تعیین مرز در گروه اسمی کمک کند اشاره شده است. مثل این که ضمیرها و اسم‌های خاص معمولاً آخرین سازه گروه اسمی هستند.

در قسمت چهارم این مقاله ساختار فرهنگ لغت شیراز معرفی شده است. این فرهنگ حدود ۵۰۰۰۰ مدخل شامل واژه‌ها، گروه‌ها و اسم‌های خاص است. صورت اسنادی^۲ در فرهنگ شیراز برای فعل، مصدر و برای عناصر غیرفعلی، ریشه یا حالت تصریف نشده است. هر ورودی به‌طور خودکار به حروف لاتین نوشته می‌شود. معادل سازی برای حروف به نحوی انجام شده است که تا بیشترین حد ممکن همه اطلاعات منتقل شود. طبیعی است که واژه‌های کوتاه که در فارسی نوشته نمی‌شوند در این سیستم هم تبدیل نمی‌شوند.

اجزای کلام موضوع دیگری است که به آن پرداخته شده است. اجزای کلام به دودسته باز و بسته تقسیم می‌شوند. در این فرهنگ لغت اسم‌های خاص، فعل‌ها و فعل‌های

¹ boundary

² Citation form

سبک به طبقه باز تعلق دارند. و طبقه بسته عناصری چون حروف، لقب‌ها و اعداد را دربر می‌گیرد.

فعل‌ها به حالت مصدری در این فرهنگ وارد شده‌اند و از آن جایی که به‌دست آوردن بن مضارع از مصدر همیشه ممکن نیست، بن مضارع فعل‌ها در یک دسته مشخص وارد شده‌اند. معنی واژه‌ها، معانی کم‌بسامد و مترادف‌ها را هم شامل می‌شود. این معانی به ترتیب از پربسامد به کم‌بسامد فهرست می‌شوند. جمع‌های مکسر عربی که در فارسی وارد شده‌اند، به صورت مستقیم و با ویژگی "جمع" وارد فرهنگ می‌شوند.

قسمت پنجم مقاله با عنوان "نحو فارسی"، با این جمله که "نحو فارسی در نوشتار بسیار مبهم است و موجب بروز مشکلات زیادی در تقطیع متن نوشتاری می‌شود"، آغاز شده است.

ترتیب آزاد کلمات، نبود حالت نمای مشخص، اختیاری بودن نهاد و نیز عدم نگارش واژه‌های کوتاه به عنوان دلایل این ابهام معرفی شده‌اند.

در ادامه به اجزای برنامه نرم‌افزار پرداخته شده است. فرآیند ترجمه انگلیسی به فارسی در این نرم‌افزار به صورت زیر تشریح شده است:

۱- خوانش و آماده‌سازی متن

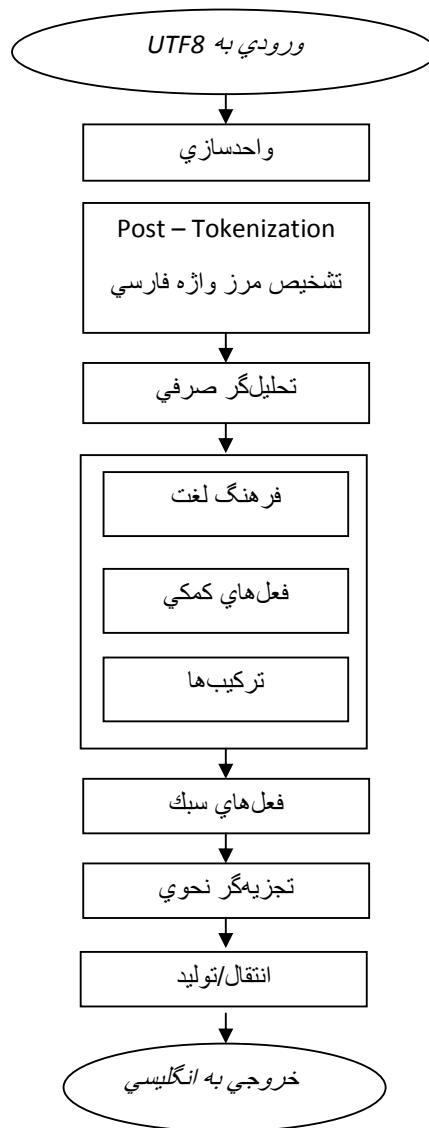
۲- تحلیل صرفی و جستجو در فرهنگ لغت

۳- تجزیه نحوی

۴- انتقال

۵- تولید و آماده‌سازی خروجی به زبان مقصد

نمودار ۱-۱



برنامه این نرم افزار با ++C نوشته شده است. نیز کمی از جاوا بهره گرفته شده است. این نرم افزار قادر به ترجمه جمله هایی با طول و پیچیدگی متوسط در مدت ۳ تا ۵ ثانیه است.

ترجمه خود کار دوسویه متون فارسی و انگلیسی^۱ مقاله ای است که سازو کار مترجم پن ترنس^۲ را بررسی می کند. پن ترنس یک مترجم متن خود کار دوسویه انگلیسی فارسی است. این مترجم شامل دو بخش اصلی است: پن تی^۳ و پن تی^۴. پن تی^۱ جمله های انگلیسی را به فارسی ترجمه می کند و پن تی^۲ مترجم فارسی به انگلیسی است. ابهام زدایی از معنای واژه که بخش مهمی از ترجمه است در پن تی یکبارگی به کار گیری ترکیبی از فرهنگ لغت گسترده^۵ و رهیافت های پیکره بنیاد، و در پن تی دو با کمک مجموعه ای از رهیافت های قاعده بنیاد، دانش بنیاد و پیکره بنیاد انجام می شود.

این مترجم دارای سه جزء اصلی است. **تحلیل ورودی^۶** که واژه های یک جمله را به همراه اطلاعات مهم درباره این واژه ها استخراج می کند. **انتقال واژه ها^۷** که هدف آن یافتن بهترین ترجمه برای هر واژه است. و **انتقال ساختاری^۸** که ترجمه استخراج شده را در ساختار زبان مقصد قرار می دهد.

¹ Automatic Translation between English and Persian Texts

² PEnTrans

³ PEnT1

⁴ PEnT2

⁵ Extended dictionary

⁶ Input analysis

⁷ Lexical transfer

⁸ Structural transfer

پن تی ۱ که مترجم انگلیسی به فارسی پن ترنس است، قادر به ترجمه جمله‌های ساده امری، منفی و سئوالی است. ظاهراً منظور از جمله‌های ساده در این نرم‌افزار جمله‌هایی است که دارای یک فعل اصلی یا حداکثر دو فعل هستند.

تحلیل ورودی: پن تی ۱ در وهله اول از سازه‌یاب ستنفورد^۱ برای واحدسازی، سازه‌یابی و بخش‌هایی از تحلیل صرفی استفاده می‌کند. پن تی ۱ از خروجی این قسمت که شامل واژه‌ها و اجزای کلام و روابط نحوی سازه‌های مختلف جمله است، استفاده می‌کند.

ابهام‌زدایی از معنای واژه زبان مبدا: در زبان مبدا (انگلیسی) از طریق گسترش خوارزمیک (الگوریتم) لسک^۲ (لسک: ۱۹۸۶) با استفاده از شبکه واژگانی^۳ (فل‌باوم^۴: ۱۹۹۸؛ بنرجی و پدرسن^۵: ۲۰۰۲) و شبکه واژگانی گسترده^۶ (هاراباگیو^۷: ۱۹۹۹).

چهارروش برای ابهام‌زدایی در این کار پیشنهاد شده است:

الف: WSD-S : خوارزمیک (الگوریتم) گسترده لسک: براساس این الگوریتم واژه‌نامه^۸ همه معانی واژه مبهم با واژه‌نامه سایر کلمات جمله قیاس می‌شود. از معانی مختلف، آن معنی‌ای انتخاب می‌شود که واژه‌نامه‌اش تعداد واژگان مشترک بیشتری با واژه‌نامه سایر

¹ Stanford syntax parser

² Lesk algorithm

³ WordNet

⁴ Felbaum

⁵ Banerjee and Pedersen

⁶ eXtended wordNet

⁷ Harabagiu

⁸ gloss

کلمات جمله دارد. در مدل به کار رفته در این نرم افزار به واژگان موجود در واژه نامه وزن های مختلف می دهند و نهایتاً به جای شمارش واژگان مشترک، مجموع وزن این واژگان محاسبه می شود.

ب: WSD- E1 : مدل الف، مشکل هم نویسه ها و چند معناها و حتی اجزای کلام مختلف را حل نمی کند. از طریق "شبکه واژگانی گسترده" برچسب های واژگانی و معنایی استخراج می شود. تنها در صورتی کاملاً مشابه تلقی می شوند که ریشه و مقوله دستوری و معنی آنها مشابه باشد.

پ: WSD- E2 : در این مدل، حرکت بر اساس اولویت صورت می گیرد که این ایده از بانرجی و پدرسون (۲۰۰۲) وام گرفته شده است. این کار با هدف فراهم کردن واژه های مرتبط و برچسب دار بیشتر، صورت می گیرد.

ت: WSD- E2 : به منظور رفع مشکل کندی الگوریتم در مورد جمله هایی که بیش از ۷ واژه مبهم دارند، از رهیافت معنایی استفاده شده است. در این رهیافت تنها واژه هایی که بر اساس خروجی تجزیه گر نحوی ستنفورد با یکدیگر رابطه دستوری دارند، با یکدیگر مقایسه می شوند.

ابهام زدایی از معنای واژه زبان مقصد: برای این قسمت از کار ترجمه یک فرهنگ لغت انگلیسی به فارسی که ترجمه دقیقی از هر معنی ارائه دهد، مورد نیاز است. از آن جایی که نه یک شبکه واژگانی فارسی در ارتباط با انگلیسی و نه ترجمه ای از شبکه

واژگانی دایره معنایی^۱ موجود بوده است، برای این کار بخشی از یک شبکه واژگانی فارسی طراحی شده است. هم‌چنین از بخش افعال پروژه فارس نت^۲ (روحی زاده ۲۰۰۸) برای دستیابی به صورت ماضی و مضارع فعل‌ها بهره برده شده است.

ابتدا معادل فارسی هر واژه از فرهنگ لغتی که نگاشت یک شبکه واژگانی است استخراج می‌شود. چنان‌چه دایره معنایی واژه بیشتر از یک واژه را شامل شود، چند مدخل شامل واژه مبهم و همسایگانش ارائه می‌شود و به گوگل سرچ^۳ فرستاده می‌شود تا طبیعی-ترین ترکیب واژه‌ها انتخاب شود.

انتقال ساختاری: در این بخش نمودار درختی جمله به فارسی به دست می‌آید.

هر نمودار شامل "گروه"^۴هایی مستقل است که می‌توانند به صورت جداگانه در جمله ترجمه شوند. نمودار درختی فارسی با رهیافت از "پایین به بالا"^۵ ساخته می‌شود. به این ترتیب که در مرحله اول ترتیب کلمات بر اساس قواعد دستوری فارسی تعیین می‌شود، سپس جایگاه گروه‌ها نسبت به هم تعیین می‌شود و نهایتاً نمودار درختی براساس قواعد نحوی زبان فارسی ساخته می‌شود.

در آخرین مرحله بر اساس نتایج حاصل از WSD ترجمه صحیح هر واژه در جای مناسب (که در بخش نمودار درختی مشخص شده است) قرار می‌گیرد.

^۱ synset

^۲ FarsNet projec

^۳ Google search engine

^۴ phrase

^۵ Bottom-up

به منظور دستیابی به یک ترجمه معقول و معنادار باید اطلاعات صرفی به دست-آمده در فاز پیش‌پردازش، به واژه‌ها اضافه شود. این مرحله که در واقع پس‌پردازش صورت می‌گیرد، برخی قواعد تصریفی^۱ و نگارشی^۲ زبان فارسی به کار گرفته می‌شود. پن‌تی ۲: مترجم فارسی به انگلیسی پن‌تی ۲ نیز دارای همان سه بخش اصلی تحلیل داده‌ها، انتقال واژه‌ها و انتقال ساختاری است.

تحلیل ورودی: پن‌تی ۲ در وهله اول از تجزیه گرنحوی کیانی (۲۰۰۹) برای واحدسازی و تحلیل تصریفی استفاده می‌کند. خروجی این مرحله یک سری ریشه واژه‌ها است که بر حسب نقش دستوری (فاعل، مفعول ...) نامگذاری شده‌اند. نیز اطلاعات صرفی این واژه‌ها و برخی برچسب‌ها (مثل اسم خاص) مشخص می‌شود.

انتقال واژه‌ها: هدف این مرحله یافتن بهترین ترجمه برای هر واژه در جمله است. ابهام بزرگترین این مرحله است.

ابهام‌زدایی از معنای واژه: در این مترجم از رهیافتی دانش‌بنیاد برای حل مسئله ابهام استفاده می‌شود. این رهیافت شامل دو مرحله است:

- ۱- استفاده از قواعد بینشی^۳: مبنای این قواعد دو فاکتور اصلی است: خود واژه و واژه‌های همسایه و میز سه ویژگی این فاکتورها یعنی نقش دستوری، مقوله دستوری و واژگانی که با آن‌ها باهم آیی دارند.

¹ inflectional
² orthographic
³ heuristic

بر این اساس برای قواعد این مرحله ۸ تعریف شده است: خود واژه، واژه‌های همسایه، مقوله دستوری و نقش دستوری واژه و واژه‌های همسایه، جایگاه^۱ واژه و نیز مقوله مفهومی واژه‌های همسایه (۲۰ مقوله مفهومی تعریف شده است). برای حل مشکل ابهام، پن تی ۲ قاعده مناسب با شرایط واژه مبهم را پیدا می‌کند.

۲- جستجو در اینترنت: چنانچه پس از مرحله اول واژه مبهمی باقی مانده باشد، آن-گاه پن تی ۲ از اینترنت به عنوان یک پیکره زبانی بزرگ استفاده می‌نماید. این مرحله پس از انتقال ساختاری (که در آن جمله کامل انگلیسی تقریباً استخراج می‌شود) انجام می‌شود. در این مرحله جمله کامل انگلیسی یا بخشی از آن در گوگل مورد جستجو قرار می‌گیرد. هر بار یکی از معانی مختلف واژه مبهم جستجو می‌شود و نهایتاً آن معنی که تعداد مثال بیشتری برایش یافت می‌شود، انتخاب می‌شود.

انتقال ساختاری: واژه‌های ترجمه شده باید بر اساس قواعد آرایش واژه‌های

انگلیسی مرتب شوند. ابتدا باید سازه^۲ های دستوری در جای مناسب قراردادده شوند. علاوه بر این به یک سری فرآیندهای دیگری هم نیاز هست تا ترجمه مطلوب‌تری حاصل شود. به این منظور بیش از ۵۰ قاعده که در ۱۳ مقوله اصلی (بر اساس ۱۳ نقش دستوری اصلی) و

¹ location

² constituent

۶ دسته جانبی جای می‌گیرند، تعریف شده است. قواعد اصلی، ترجمه هر نقش دستوری را در جایگاه خود قرار می‌دهند و قواعد جانبی یا برای آماده کردن گروه‌های فارسی برای ترجمه به انگلیسی به کار می‌روند، یا برای بهبود جمله انگلیسی تولید شده.

مقاله "ابزارهای اساسی برای پردازش متون فارسی"^۱ این مقاله ابزارهایی کلیدی

برای پردازش متون فارسی را معرفی می‌کند. سامانه معرفی شده در این مقاله Step-1 نام دارد و قادر است متن فارسی دریافتی را از نظر صرفی تحلیل، و اجزای کلام را نشان‌گذاری کند.

مقاله "پردازش متون فارسی: دستاوردهای گذشته، چالش‌های پیش‌رو". مهنوش

شمس فرد. (۱۳۸۵). این مقاله به بررسی پردازش متون فارسی و مسائل مطرح در آن می‌پردازد. در این راستا ابتدا به بررسی و تحلیل فعالیت‌های انجام شده در رابطه با پردازش متون فارسی در ایران و سایر کشورهای جهان در سال‌های اخیر می‌پردازد و سپس با طرح مسائل و چالش‌های این امر نگاهی به فعالیت‌های مورد نیاز آتی برای رساندن تحقیقات گذشته به کاربردهای آینده دارد.

در مقاله "صرف زبان فارسی از دیدگاه رایانه"^۲ یک بررسی از صرف زبان فارسی

با دیدگاه کاربرد در رایانه ارائه می‌شود. اجزای کلام و تکواژها نیز در این مقاله بررسی

¹ A Set of Fundamental Tools for Persian Text Processing

² Persian Computational Morphology

می‌شوند. هم‌چنین به پارادایم فعلی نیز پرداخته می‌شود. در این گزارش تحلیل‌گر صرفی پروژه شیراز هم بررسی می‌شود و نمونه‌هایی از قواعد آن نیز ارائه می‌شود.

هشام فیلی (۱۳۸۳) در مقاله‌ای با عنوان "ترجمه ماشینی انگلیسی به فارسی بر مبنای نظریه احتمالات و دستور کلمه" نمونه‌ای از یک سامانه مترجم ماشینی انگلیسی به فارسی

مبتنی بر دستور درخت افزایشی هم زمان ارائه می‌دهد. در نمونه معرفی شده ترجمه صرفاً از انگلیسی به فارسی است، ولی نویسنده معتقد است با وجود ویژگی‌های این مدل امکان به‌کارگیری آن در جهت عکس نیز وجود دارد. اولین مرحله از ترجمه در این روش "فاز تجزیه‌گر" نامیده شده که در آن جمله در زبان مبدا مورد تجزیه قرار می‌گیرد. مرحله بعدی "فاز انتقال" است که خود شامل سه بخش است. در این فاز یک درخت اشتقاق زبان مبدا را به عنوان ورودی گرفته و یک درخت اشتقاق زبان مقصد را خروجی می‌دهد و در واقع هسته اصلی جمله ترجمه می‌شود. سه بخش این فاز عبارتند از: انتقال درختی، انتقال واژگانی و انتقال ویژگی. سومین تحلیل ساختاری است. در ادامه ترجمه چند نمونه ارائه شده است. در این مقاله دستور درخت افزایشی هم زمان به عنوان روشی کارآمد معرفی شده که نقایص خود را هم دارد. از جمله این که میزان استفاده از اطلاعات معنایی و مفهومی در این مدل ناکامل است، و در نتیجه گاهی قادر به رفع ابهامات واژگانی نیست

تعداد زیادی مقاله وجود دارد که می‌توان هریک را به نحوی با موضوع ترجمه

ماشینی مرتبط دانست. از جمله:

- ترجمه ماشینی انگلیسی به فارسی بر مبنای نظریه احتمالات و گرامر لغوی (فیلی: ۱۳۸۳)

- پیکره متنی زبان فارسی (بیجن خان و مرادزاده: ۱۳۸۳)

- طراحی یک واژگان جامع برای پردازشگر زبان فارسی (نوربالا: ۱۳۸۰)

هم چنین فهرستی از پایان‌نامه‌های ارائه‌شده در زمینه "زبان فارسی و رایانه" از ۱۳۵۵-

۱۳۸۴ در مقاله‌ای با همین عنوان (ناصح: ۱۳۸۵) در مجموعه مقالات دومین کارگاه

پژوهشی زبان فارسی و رایانه ارائه شده است.

بدیهی است ارائه تمام مقالات و پژوهش‌هایی که ممکن است به نحوی با موضوع

ترجمه ماشینی مرتبط باشند غیرضروری به نظر می‌رسد. چرا که ممکن است به جای

کمک به روشن‌تر شدن موضوع موجب سردرگمی شود. تکرار مفاهیم نیز در بسیاری

از منابع، دلیل دیگری برای این امر است.

به هر روی تلاش نگارنده براین بوده است که ضمن بررسی منابعی که دسترسی به

آنها ممکن بوده، بتواند تصویر مناسبی از روند پیدایش و پیشرفت ترجمه ماشینی در

ایران و جهان ارائه دهد. در مواردی که توضیح‌های بیشتر ضروری به نظر می‌رسید

مقاله‌ها به تفصیل مورد بررسی قرار گرفتند، در غیر این صورت به ارائه توضیح

مختصری بسنده شد.

فصل دوم: ترجمه ماشینی در زبان فارسی

چنان که در بخش اول هم توضیح داده شد ترجمه ماشینی یک پژوهش میان رشته-ای است که در آن حداقل به دو تخصص نیاز است. در اغلب منابع مرتبط به این حوزه بر جنبه برنامه نویسی ترجمه ماشینی تأکید شده و غالباً تحقیقات زبان‌شناختی لازم برای ساختن مترجم ماشینی را مفروض انگاشته‌اند.

در فصل حاضر قصد داریم به ارائه طرح‌های مرتبط با تحقق ترجمه ماشینی در زبان فارسی بپردازیم. واقعیت این است که مواد لازم برای ساختن مترجم ماشینی از زبان مورد تحلیل به دست می‌آید. چنان‌چه زبان مفروض مورد تحلیل دقیق و متقن قرار نگیرد نمی‌تواند براساس آن برنامه رایانه‌ای منسجم و کاملی را عرضه نمود. ساختار زبان مورد تحلیل باید به دقت استخراج و به خصوص در حوزه‌های واژگان و نحو آماده برنامه نویسی شود. یعنی واژگان لازم بسته به سبک مورد تحلیل از پیکره زبان فارسی انتخاب و با ویژگی‌های صرفی و نحوی خود در واژگان قرار می‌گیرد. این واژگان کاربردی نظیر واژگان ذهنی اهل زبان خواهد داشت. در مرحله بعدی باید سازه‌های زبان مورد تحلیل، و نحو آن به دقت و صراحت مشخص شود و به شکل زبانی صوری و آماده برنامه نویسی درآید.

۲-۲- زبان فارسی

زبان فارسی به خانواده زبان‌های هندواروپایی تعلق دارد. تعداد سخنگویان زبان فارسی از سخنگویان دیگر زبان‌های جدید ایرانی بیشتر است (ماهوتیان: ۱۳۷۸).

۲-۲-۱- خط فارسی

زبان فارسی با تغییراتی در خط عربی و از راست به چپ نوشته می‌شود و سی‌و‌دو حرف دارد.

از جمله مشکلات خط فارسی می‌توان به موارد زیر اشاره کرد:

- وجود چند نویسه برای یک آوا (س، ث، ص برای /s/، و ت، ط برای /t/،
- وجود یک نویسه برای چند آوا (ـُ، و برای /o/،
- چند شکلی بودن حروف با توجه جای قرار گرفتن در کلمه (مثل ص در ابتدای کلمه، ص در میان کلمه، ص چسبیده به انتهای کلمه، ص در انتهای کلمه و جدا) زیرا خط فارسی "پیوسته نویس" است و در نگارش آن حروف به هم می‌چسبند.

- عدم نگارش مصوت‌های کوتاه (ـَ، ـِ، ـُ)

از دیدگاه ترجمه ماشینی:

- دشواری در تعیین مرز کلمه (بویژه در مورد کلمات پیچیده مثل کلمات مرکب و

اشتقاقی)

۲-۳- ترجمه متن فارسی

به طور کلی برای ترجمه متن، ابتدا باید آن متن در زبان مبدا مورد تحلیل قرار بگیرد. برای ساختن یک نرم افزار ترجمه به سه بخش متفاوت نیاز است: یک بخش واژگان، یک تحلیل گرسرفی و یک تجزیه گرنحوی. انتخاب واژه مناسب و ترجمه جمله با ساختاری صحیح در زبان مقصد، ترجمه نهایی جمله را به دست می دهد. در واقع پس از تحلیل متن در زبان مبدا، امکان انتقال آن به زبان مقصد میسر می شود. این در حالی است که چنان تحلیلی در زبان مقصد هم امکان پذیر است (به این معنی که زبان مقصد هم، به همان میزان مورد تحلیل و مطالعه قرار گرفته و قواعد آن استخراج شده است^۱).

۲-۳-۱- استخراج واژه ها

قبل از جستجوی واژه در فرهنگ، نرم افزار باید قادر به تشخیص آن باشد. یعنی باید بتواند مرز کلمه را تعیین کند.

^۱ تجزیه و تحلیل صرفی و نحوی باید در هر دو زبان مبدا و مقصد انجام شود. در اینجا متناسب با موضوع طرح، مثال ها در

مورد زبان فارسی بررسی می شود.

یکی از راهکارهای ابهام‌زدایی در تعیین مرز کلمه، استفاده صحیح از "فاصله" و "نیم‌فاصله" در نوشتن متن است. اگر بین واژه‌ها از "فاصله" کامل، و بین اجزای تشکیل دهنده یک کلمه پیچیده (در صورتی که جدانوشته شود، مثل "می‌روم") از "نیم‌فاصله" استفاده شود (مثل "نرم‌افزار" به جای "نرم افزار")، نرم‌افزار قادر به تشخیص مرز کلمه خواهد بود.

به هر حال تحلیل‌گر صرفی نرم‌افزار نیز در تعیین مرز کلمه نقش مهمی بازی خواهد کرد. با توضیح بخش‌های ضروری این تحلیل‌گر چگونگی این تاثیرگذاری آشکار خواهد شد.

۲-۳-۲- واژگان^۱

در واژگان اطلاعات آوایی، صرفی، نحوی و معنایی یک کلمه باید گنجانده شود. اطلاعات آوایی شامل نحوه تلفظ واژه است با استفاده از علائم آوانگار بین‌المللی باید در واژگان نگاشته شود. اطلاعات صرفی و نحوی به مقولات دستوری مانند اسم و فعل و ... اشاره دارد. معانی متعدد یک واژه باید بر اساس بسامد کاربرد آن، در واژگان گنجانده شود.

¹ lexicon

۲-۳-۳- تحلیل صرفی و نحوی

تحلیلی نظام‌مند از ساختار زبان را دستور می‌نامند. دستور زبان را از یک دیدگاه می‌توان به دستور تجویزی و توصیفی تقسیم کرد. دستور توصیفی دستوری است که در آن به توصیف زبان آن‌چنان که هست، دست می‌زنند. دستور تجویزی شامل قواعدی برای کاربرد درست زبان در جامعه است (کریستال: ۱۹۹۲).

آن‌چه که در تحقیقات زبانی در حوزه پردازش زبان کاربرد دارد باید توصیفی باشد. یعنی توصیف صوری زبان آن‌چنان که هست، نه آن‌چنان که باید باشد.

بر اساس تعریف سنتی دستور در زبان‌شناسی، این واژه به مطالعه ساختاری زبان

فارغ از آواشناسی^۱ و معنی‌شناسی^۲ دلالت می‌کند و به دو بخش "صرف"^۳ و "نحو"^۴

تقسیم می‌شود (کریستال: ۱۹۹۲).

به طور کلی می‌توان گفت که صرف به بررسی اجزای کلمه می‌پردازد و نحو به بررسی ساختار گروه و جمله.

۲-۳-۳-۱- صرف

ذخیره همه واژه‌های یک زبان در حافظه نرم‌افزار اقتصادی به نظر نمی‌رسد. با کمک صرف می‌توان با ادغام اجزا و عناصر موجود در زبان کلمات جدید را تولید کرد.

^۱ phonology

^۲ semantics

^۳ morphology

^۴ syntax

صرف شامل دو بخش "تصرف" و "واژه سازی" است. تصرف افزودن وند به کلمه است به نحوی که مقوله نحوی واژه نهایی با واژه پایه یکی است. مثل افزودن "ها" جمع به "اسم".

واژه سازی شامل اشتقاق^۱ و ترکیب^۲ است. کلمات مشتق از یک پایه^۳ به اضافه وند ساخته می شوند. مانند دانشمند، بی باک. کلمات مرکب از اجتماع اسم و صفت و قید و بن فعل با هم ساخته می شوند، نظیر کتابخانه، عقب گرد، خانه نشین.

با توجه به آنچه گفته شد یک تحلیل گر صرفی بخش مهمی از نرم افزار مترجم را تشکیل می دهد. برای تهیه چنین تحلیل گری به یک مجموعه قواعد مدون نیاز است که با دیدگاه کاربرد در ترجمه ماشینی تدوین شده باشد. تحلیل گر مفروض باید قادر به شناسایی کلمات و اجزای تشکیل دهنده آنها باشد. به این ترتیب برچسب گذاری اجزای کلام نیز میسر می شود.

برای توضیح بیشتر آنچه که در تدوین قواعد مورد نیاز است، بهتر دیدیم به معرفی دو نرم افزار واژه شکن و فعل شکن پردازیم که نگارنده با هدف به کارگیری در ترجمه ماشینی اقدام به تهیه آن نموده است.

۲-۳-۱-۱- واژه شکن و فعل شکن: نرم افزارهایی برای تحلیل صرفی

^۱ derivation

^۲ compounding

^۳ base

۲-۳-۱-۱-۱- واژه شکن

واژه شکن فارسی برای تحلیل کلمات پیچیده ساخته شده است. این نرم افزار کلمات مشتق

و مرکب را تجزیه و مقوله آن ها را مشخص می کند.

واژه شکن فارسی اهداف زیر را تامین می کند:

الف. تجزیه کلمات مشتق به ریشه و وند (پسوند و پیشوند)

ب. تجزیه کلمات مرکب به عناصر سازنده آن

ج. تعیین مقوله دستوری کلمات مشتق و کلمات پیچیده

ساختمان واژه شکن

واژه شکن در بخش عمده دارد: بخش الفاظ و بخش قواعد. بخش الفاظ شامل انواع

مقولات دستوری (اسم، صفت، عدد، قید، ضمیر، بن فعل) است. وندها (پیشوندها و

پسوندها) را هم باید جزء این بخش دانست.

الف. بخش اسم ها

در این بخش سی و پنج نوع اسم وجود دارد: اسم، اسم جاندار، اسم معنی، اسم عمل، اسم

زمان، اسم خویش، اسم مصدر، اسم مرکب، اسم جاندار مرکب، اسم معنی اشتقاقی، اسم

بن ساختی، اسم جاندار بن ساختی، اسم خویش مرکب، اسم جاندار اشتقاقی، اسم عمل

اشتقاقی، اسم اشتقاقی، اسم خویش اشتقاقی، اسم ظرف اشتقاقی، اسم مکان اشتقاقی، اسم

مجموعه، اسم زمان اشتقاقی، اسم خویش مرکب بن ساختی، اسم معنی مرکب اشتقاقی،
اسم معنی بن ساختی اشتقاقی، اسم مجموعه مرکب، اسم مجموعه بن ساختی، اسم خویش
مرکب اشتقاقی، اسم مرکب بن ساختی، اسم مکان مرکب اشتقاقی، اسم مکان بن ساختی
اشتقاقی، اسم مرکب اشتقاقی، اسم بن ساختی اشتقاقی، اسم جاندار مرکب اشتقاقی، اسم
جاندار بنساختی اشتقاقی، اسم عمل مرکب اشتقاقی، اسم عمل بن ساختی اشتقاقی.

ب. بخش صفت ها

در فهرست صفت ها دوازده نوع صفت تشخیص داده شده است: صفت، صفت مفعولی،
صفت مفعولی منفی، صفت مبالغه، صفت مبالغه منفی، صفت مرکب، صفت بن ساختی،
صفت ترتیبی، صفت اشتقاقی، صفت مرکب اشتقاقی، صفت بن ساختی اشتقاقی، صفت
مرکب بن ساختی.

ج. بخش بن ها

شامل آخرین بن مضارع و بن ماضی افعال. آن دسته از بن های مضارع و ماضی که تبدیل
به اسم شده اند و یا ابتدا اسم بوده اند و به صورت مصدر جعلی درآمده اند در فهرست بن
ها واسامی مشترک اند و در هر دو جا درج شده اند.

د. بخش قیود و ضمائر

در این فهرست دونوع قید (قید، قید جهت) و یک نوع ضمیر وجود دارد.

ه. بخش اعداد

در این طبقه اعداد یک تا بیست و هم چنین سی تا صد (به صورت متناوب ده تایی) وجود دارد

بخش وندها

شامل شش پیشوند (مبالغه، دارندگی، منفی ساز، منفی ساز ۲، اشتراک).
سی و هشت پسوند نیز در فهرست وندها آمده است: شباهت، آغستگی، اسم معنی مرکب ساز، مبالغه، تقلید، محافظت، ظرف، شباهت ۱، شباهت ۲، شباهت ۳، شباهت ۴، شباهت ۵، شباهت ۶، تصغیر، تصغیر ۱، نسبت، نسبت ۱، مکان، مکان ۱، مکان ۲، مکان ۳، مجموعه، کنندگی، کنندگی ۱، عمل، عمل ۱، عمل ۲، دارندگی، دارندگی ۱، دارندگی ۲، دارندگی ۳، ترتیب، ترتیب ۱، جنس، فاعلی، فاعلی ۱، مصدر ساز، مفعولی.

بخش قواعد

در بخش فهرست قواعد انواع امکانات هم نشینی تکواژها و وندها با یکدیگر به صورت قاعده آمده است. در این جا شیوه استخراج قواعد به اختصار ذکر می شود.

نمونه‌ای از قواعدواژه‌شکن:

کلمات اشتقاقی
اسم + وند شباهت = صفت اشتقاقی
اسم مرکب + وند شباهت = صفت مرکب اشتقاقی
کلمات مرکب

اسم + صفت = اسم جاندار مرکب
اسم اشتقاقی + صفت = اسم جاندار
کلمات بن ساختی
اسم + بن مضارع = اسم جاندار بن ساختی
اسم اشتقاقی + بن مضارع = اسم بن ساختی اشتقاقی

شیوه استخراج قواعد

واژه شکن همانند مغز انسان عمل می کند. یعنی به جای تکیه کردن بر محفوظات از قواعدی حتی المقدور حشودایی شده و تعمیم پذیر استفاده می کند. هر تکواژ (چه تکواژ مستقل و چه مقید) برای قرار گرفتن در کنار تکواژهای دیگر محدودیت هایی دارد. بدین معنی که قبل و بعد از بعضی تکواژها می آید و در کنار برخی دیگر نمی آید. کلمات پیچیده نیز به نوبه خود تابع محدودیت هایی هستند. برای دست یافتن به امکانات و محدودیت های هم نشینی تکواژها و کلمات پیچیده ابتداء کلمات پیچیده شامل ترکیبات اشتقاقی و غیر اشتقاقی و کلمات مشتق مرکب و ترکیبات بن ساختی از فرهنگ معین استخراج شد. علت انتخاب فرهنگ معین برای استخراج داده ها این بود که ساخت های قدیمی و و مهجور نیز در این فرهنگ وجود دارد. در مرحله بعد فهرستی از اجزای تشکیل دهنده ترکیبات استخراج شده تهیه شد. یعنی مقوله های متعارف دستوری نظیر اسم و صفت و قید و ضمیر و غیره و هم چنین وندها (پیشوندها و پسوندها) که در ساختمان ترکیبات اشتقاقی و غیر اشتقاقی و بن ساختی و مشتق مرکب به کار رفته بودند جدا و طبقه

بندی شدند. داده‌های استخراج شده اولیه در طی کاردچار جرح و تعدیل های فراوانی شدند که به مقتضای کار بود.

اجزاء استخراج شده از ترکیبات پیچیده یک به یک در بافتی که به کار رفته بودند و بافت‌های بالقوه‌ای احتمال بکار رفتن آن‌ها بود بررسی شدند و به هریک برچسبی داده شد. در این مرحله مشخص شد که مقوله‌های متعارفی که محققان تا کنون در آثارشان تشخیص داده و نامگذاری کرده اند کفایت نمی کند و باید از مولفه های مکمل برای نامیدن اجزاء آن استفاده کرد. در تحقیقاتی که تا کنون درباره مقوله های دستوری شده است از مقوله های اسم و صفت و ضمیر و قید و حرف اضافه و حرف ربط و اصوات نام برد اند. محمودی (۱۹۹۴) از جمله محققانی است که به پانزده مقوله دستوری در زبان فارسی قایل شده است. آن چه در کار واژه شکن اهمیت داشت تعداد کل مقوله ها نبود بلکه مقوله های فرعی هر مقوله بود. بدین معنی که مقوله اسم بسته به نقشی که در هر کلمه پیچیده داشت باید مولفه ای جداگانه می گرفت و برچسب جدیدی به آن نسبت داده می شد. محمودی عقیده دارد که مقوله های دستوری باید بر اساس توزیع تعیین شوند (همانجا). استخراج مولفه های هر مقوله بر اساس همین نظر محمودی انجام نشد. بدین ترتیب برای هر مقوله به مقوله هایی معنایی در چارچوب مقوله اصلی قائل شدیم و بر اساس آن ها مقوله ها به چند مقوله جدید تقسیم شدند. سی و پنج نوع اسم و دوازده نوع صفت حاصل همین کار بودند.

برای قاعده نویسی دو کار انجام شد:

الف. قاعده هم نشینی مقوله ها به زبان طبیعی نوشته شد.

ب. قاعده ها به صورت حروف قراردادی لاتین درآمدند تا مناسب برنامه نویسی رایانه ای

شوند. مثلاً از حروف s و n و a و غیره و اندیس های عددی برای نشان دادن قواعد

استفاده شد.

نمونه:

Formula Descriptions	
Code	Description
A01	صفت
A02	صفت مفعولی
A03	صفت مفعولی منفی
A04	صفت مبالغه
A05	صفت مبالغه منفی
A06	صفت مرکب
A07	صفت بن ساختی

Formula	
Formula	Small Formula
A01I01	N12
A01I02	A07
A01N01	N09
A01N02	N09
A01N03	N09
A01N04	N09
A01N05	N09
A01N07	N09

ویژگی های واژه شکن

یکی از مشکلات نظری در تهیه واژه شکن مسئله دستوری و غیر دستوری بودن ساخت ها

بود. غرض از ساختن واژه شکن آن بود که نرم افزار مفروض بتواند ساخت های دستوری

را از ساخت های غیردستوری را تمیز دهد و سپس الگوی ساخت آن را عرضه کند.

بنابراین مشکل اصلی این بود که واژه شکن چه محدوده ای را مجاز و چه محدوده ای را غیر مجاز بداند.

سؤال این بود که اگر قرار باشد ساختی را غیردستوری بدانیم معیار ما چه خواهد بود. چامسکی در ساخت های نحوی آن جا که درباره استقلال دستور زبان بحث می کند به این نکته می پردازد که چگونه جمله دستوری را از جمله غیردستوری تمیز دهیم (سمیعی، ۱۳۶۲، ص ۱۴ تا ۲۰): در رشته های دستوری و غیر دستوری را در عمل بر چه پایه هایی از یکدیگر باز می شناسیم؟ " پاسخ های چندی که بلافاصله به ذهن خطور می کنند نمی توانند درست باشند ... دستور زبان پرتو رفتار گوینده است که بر پایه تجربه محدود و تصادفی خود از زبان می تواند تعداد نامحدودی جمله تازه تولید کند.... در حقیقت هر توضیحی درباره مفهوم دستوری بودن در زبان L را می توان بدست دادن توضیح این جنبه اساسی از رفتار زبانی شمرد. " (همانجا ص ۱۵). به نظر چامسکی " مفهوم دستوری بودن را با مفهوم بهره ور بودن یا معنی دار بودن.... نمی توان همسان دانست. " (همانجا، ص ۱۵)

او جستجو " برای یافتن تعریف دستوری بودن بر پایه معنی شناسی " را کاری عبث می داند (همانجا، ص ۱۵). از طرفی دستوری بودن را تابع بسامد و " از لحاظ تخمین آماری در حد بالا بودن " نیز نمی داند (همانجا، ص ۱۶).^۱

^۱ نگارنده این گزارش در رساله دکتری خود در پی طراحی ماشین مفروضی بوده که جمله های دستوری را بدون توسل به معنی از جمله های غیردستور تمیز دهد. یعنی بر مبنای نظریه چامسکی در کتاب ساخت های نحوی.

سه نکته مهم در نقل قول های مذکور وجود دارد:

الف. رفتار زبانی یکی از ملاک های تمیز جمله دستوری از غیر دستوری است.

ب. دستوری بودن را می توان فارغ از معنی بررسی کرد.

ج. بسامد وقوع تأثیری در دستوری یا غیردستوری بودن ندارد.

این بحث گرچه در حوزه نحو صوت گرفته است تعمیم آن به حوزه صرف هم امکان

پذیر است و می توان نتیجه گرفت که:

می شود در مورد معنی ترکیبات پیچیده تا حدی با تسامح برخورد کرد و غیرمتعارف

بودن معنی آن ها را دلیل غیردستوری بودن ترکیبات ندانست.

واژه شکن با این دیدگاه ساخته شده است. یکی از تبعات برگزیدن چنین

دیدگاهی این است که بسیاری از ساخت های بالقوه را که خوش ساخت¹ هستند ولی چه

بسا تا کنون کسی به کار نبرده است می توان دستوری دانست. دستوری دانستن ترکیباتی

نظیر چوب نشین و گرفتارانه که به قیاس با خانه نشین و دانشمندانه ساخته شده اند به

همین علت است.

¹ well- formed

بی‌نظمی موجود در ساخت‌های پیچیده (به خصوص بی‌نظمی ترکیبات اشتقاقی) نیز دلیل موجهی برای دستوری دانستن ترکیبات پیچیده بالقوه‌ای است که معنی غیرمتعارف دارند مثلاً پسوند "۱-ا" را می‌توان به صفت‌های "پهن" و "دراز" اضافه کرد و "پهنا" و "درازا" را ساخت. چنانچه این قاعده را تعمیم ندهیم و ساخت‌هایی نظیر "نازکا" و "سختا" و "نرما" را دستوری ندانیم ناچار خواهیم شد به استثناءهای زیادی قائل شویم که خلاف نظریه‌های یادگیری زبان است. یعنی مجبور خواهیم شد همه این قاعده‌ها را تعمیم ناپذیر فرض کنیم و ترکیبات موجودی را که با این قاعده ساخته شده‌اند مستقیماً در حافظه رایانه قرار دهیم. هرچند در مواردی نیز چاره‌ای جز قراردادن مستقیم ترکیبات پیچیده در حافظه وجود نداشته است. مثلاً ترکیباتی که با بن‌ماضی و پسوند "۱-ار" ساخته می‌شوند چنان از لحاظ معنایی ناهمگون‌اند که نمی‌توان برای آن‌ها قاعده عام داد و باید آن‌ها را مستقیم در حافظه گذاشت.

یکی دیگر از مشکلات نظری این بود که آیا قواعد باید دوری^۱ باشد یا خیر.

قواعدی را که می‌توانند تکرار شوند و جمله بسازند قواعد دوری یا مکرر^۲ می‌نامند (کریستال، ۱۹۹۲، ص ۳۲۸). مثلاً قاعده قرار گرفتن صفت بعد از اسم در زبان فارسی قاعده‌ای دوری است زیرا منطقاً می‌توان گفت: این کتاب قطور خواندنی جذب

سبز.....

^۱ recursive

^۲ iterative

این مسئله در کار واژه شکن که در حوزه صرف عمل می کند نیز صدق می کند.

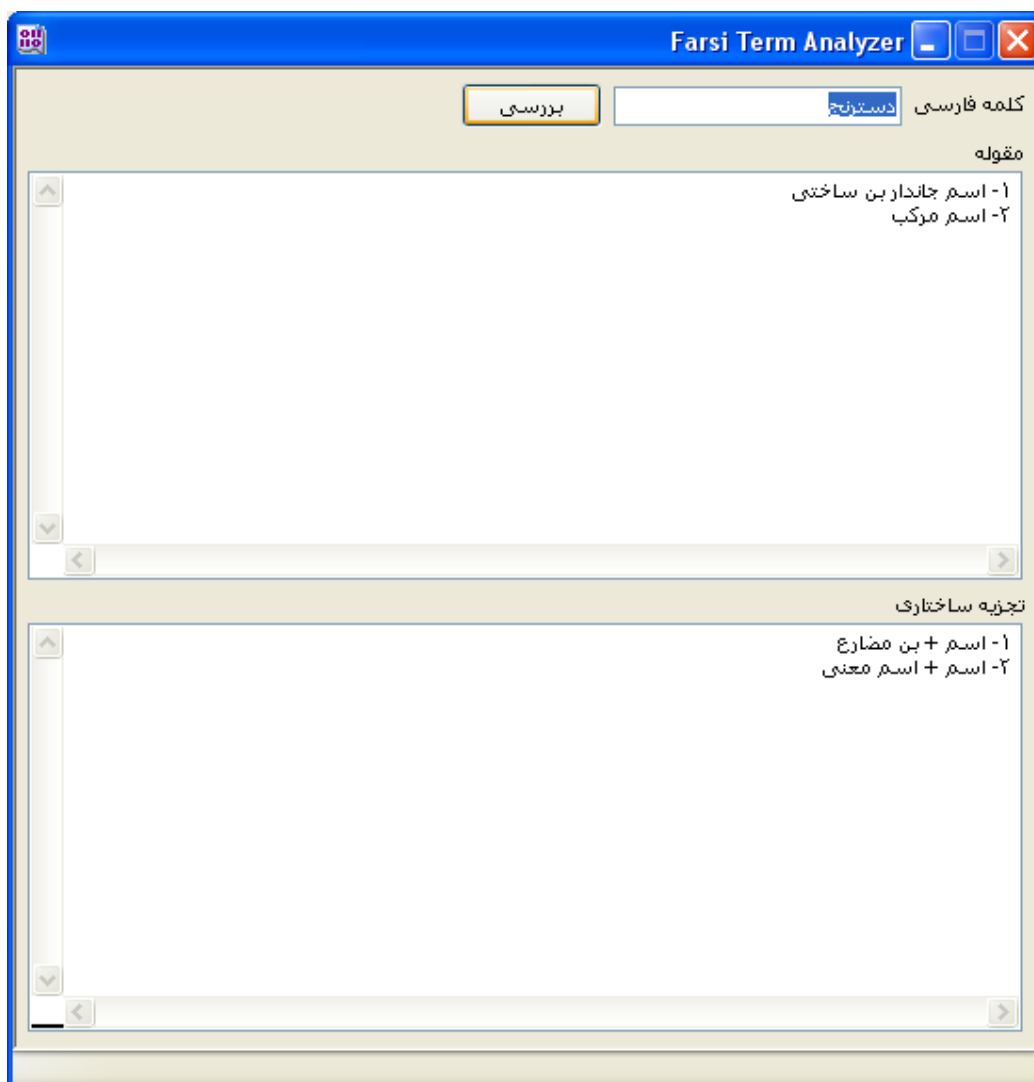
مثلا اسم + وند تصغیر (کتاب + ک) ← کتابک، که کتابک خود اسم است و منطقاً بار دیگر با پسوند " - ک" ترکیب می شود و کتابکک را می سازد و الی آخر. یا اسم جاندار + وند مبالغه (کبوتر باز) اسم جاندار که " کبوتر باز" خود اسم جاندار است و منطقاً بار دیگر با پسوند " - باز" ترکیب می شود و بار دیگر "کبوتر باز باز" را می سازد و قس‌الهدا. در زبان عادی کاربرد مانع از دوری شدن ساخت‌ها می شود. در کار واژه شکن مشکل به این صورت حل شد که برای کتابک و کبوتر باز نامی به جز اسم و اسم جاندار داده شد (اسم تصغیر اشتقاقی و اسم جاندار مرکب) و قاعده دیگری برای ترکیب این ساختهای جدید با پسوند " - ک" و " - باز" داده نشد تا مانع از دوری شدن قواعد شود.

قائل شدن به دوری بودن قواعد در کار واژه شکن باعث ساخته شدن ترکیباتی دستوری می شود. اما از آنجایی که این ترکیبات انتهای نداشتند و در کارائی واقعی واژه شکن اخلال به وجود می آورند، این ویژگی مهار شده است. بدین معنی که قواعد تکرارپذیر نیستند و در مرحله ای متوقف می شوند.

شیوه تجزیه ترکیبات و تعیین مقوله دستوری

واژه شکن ابتدا مقوله کلمه پیشنهادی را مشخص و سپس آن را به اجزاء سازنده اش تجزیه می کند. در صورتی که قاعده ساخت کلمه پیشنهادی در حافظه رایانه موجود نباشد فقط به

تجزیه کلمه مذکور اکتفا می شود. همه این موارد مشروط به این معنی است که عناصر سازنده کلمه مفروض در حافظه موجود باشد. در صورت موجود نبودن کلمات پیشنهادی در حافظه واژه شکن می توان آن کلمه را به فهرست افزود و سپس آن را برای تحلیل به واژه شکن سپرد. مثال:





 **Farsi Term Analyzer** [min] [max] [close]

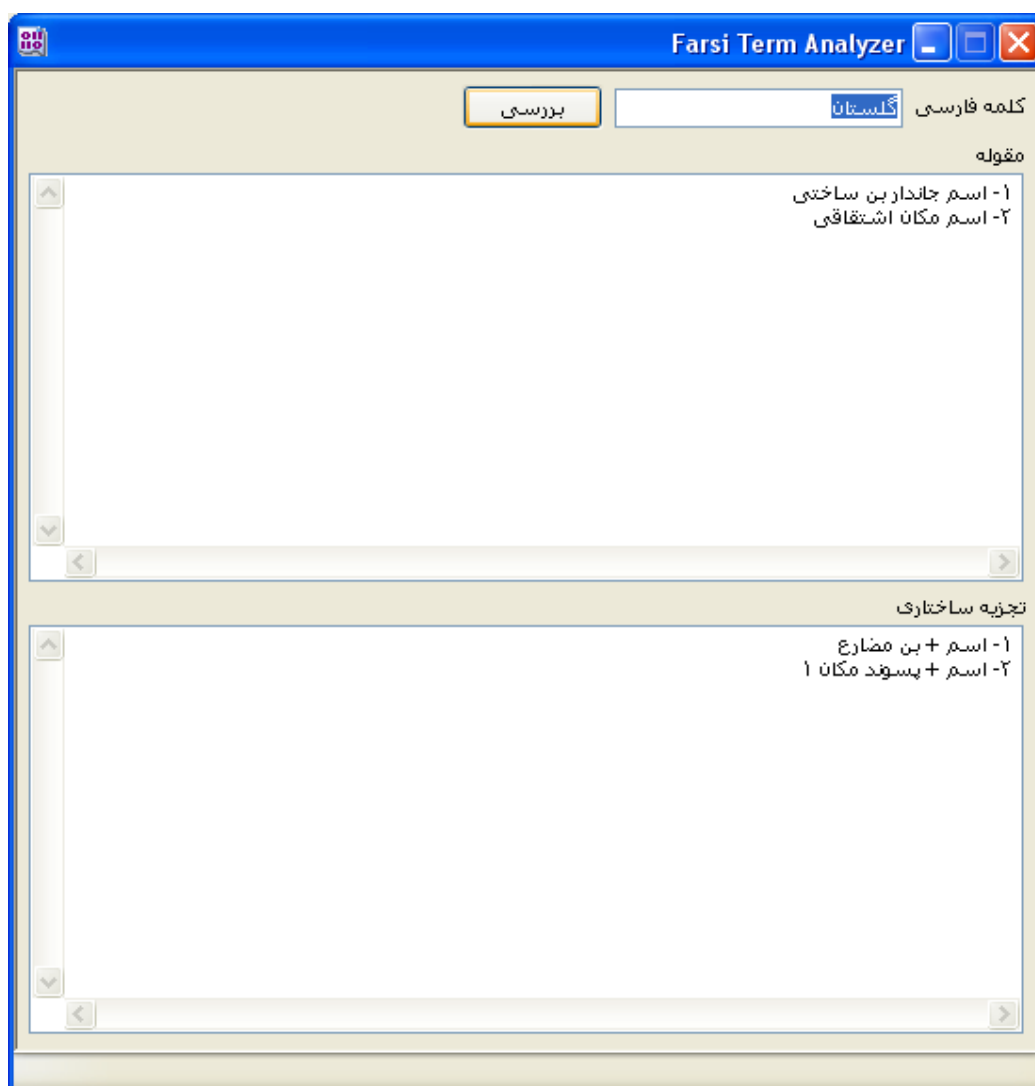
کلمه فارسی

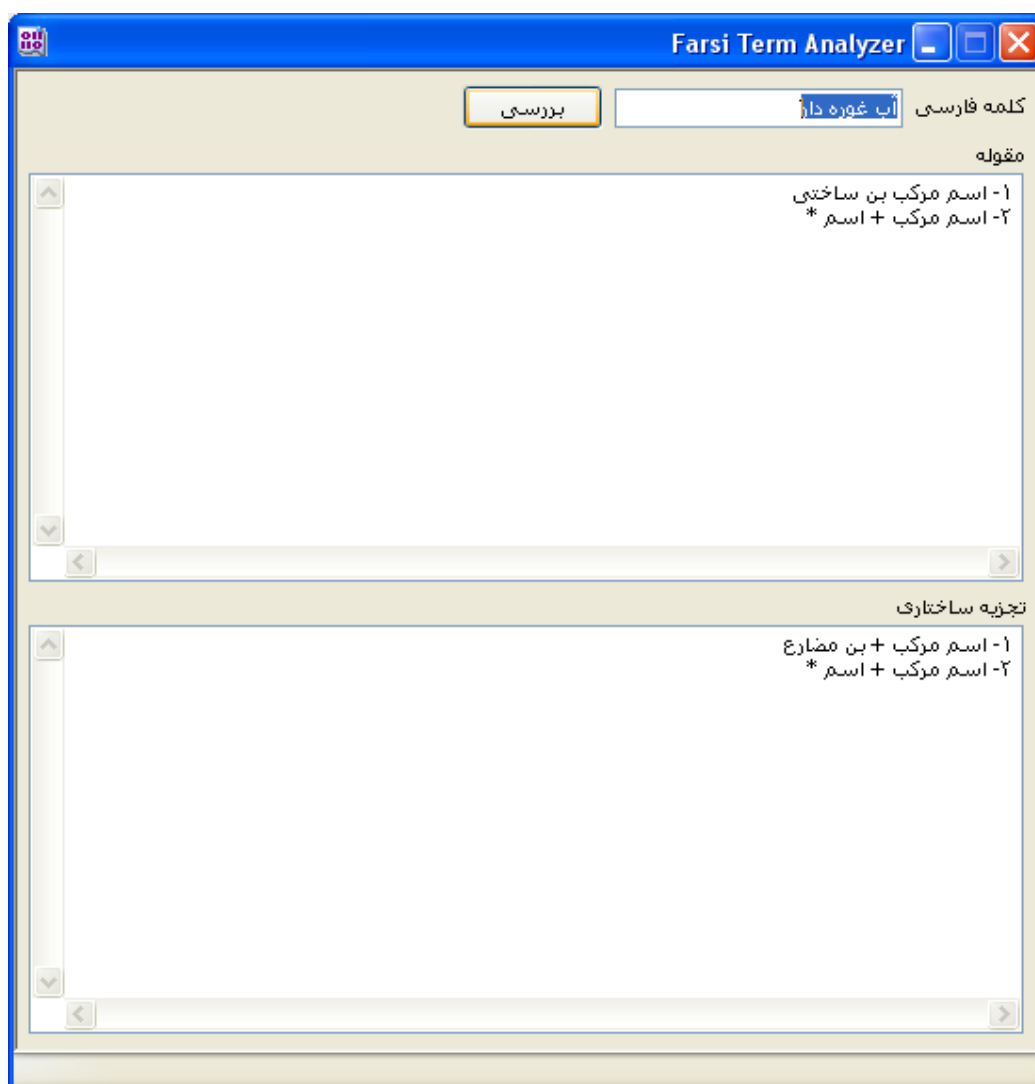
مقوله

۱- اسم جاندار اشتقاقی
۲- اسم مصدر + پسوند دارندگی ۱ *

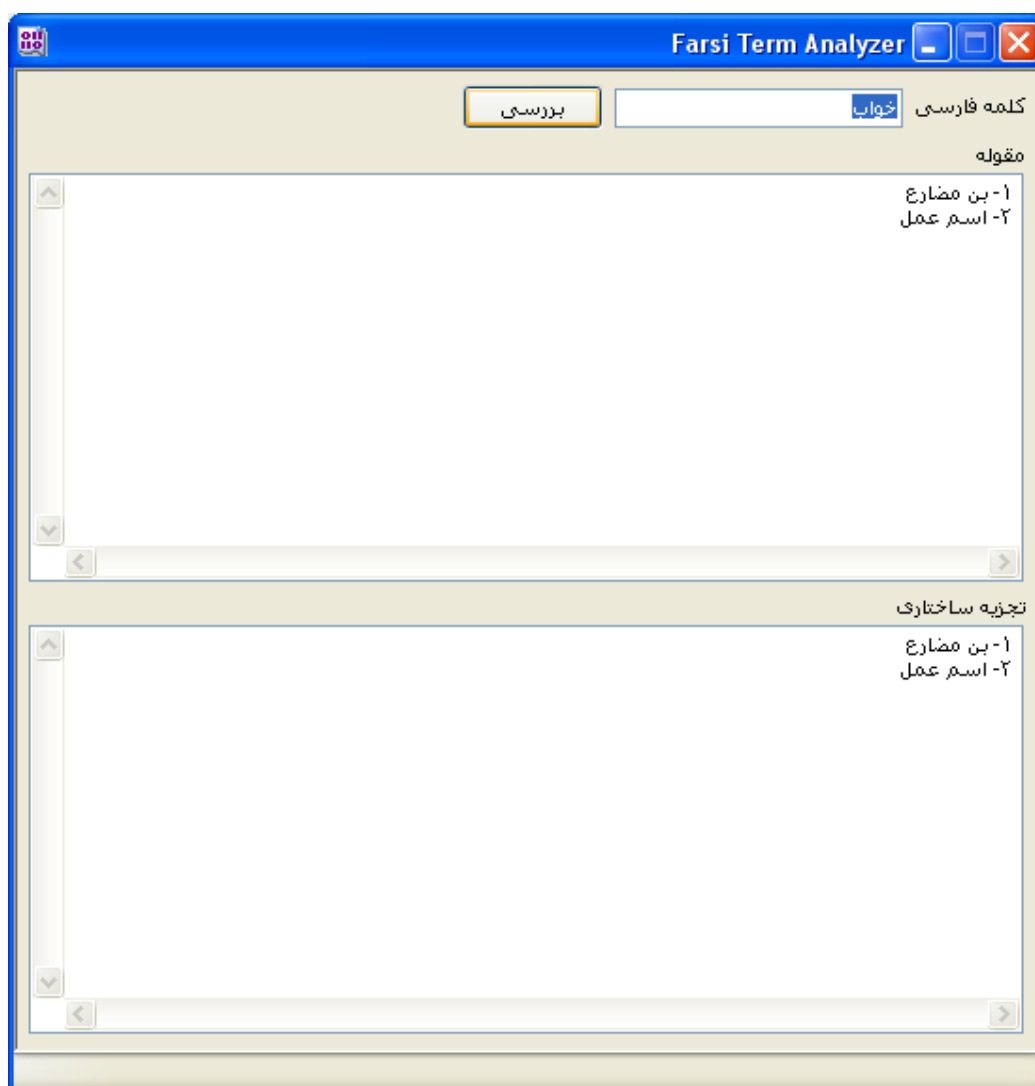
تجزیه ساختاری

۱- اسم معنی + پسوند دارندگی ۱
۲- اسم مصدر + پسوند دارندگی ۱ *

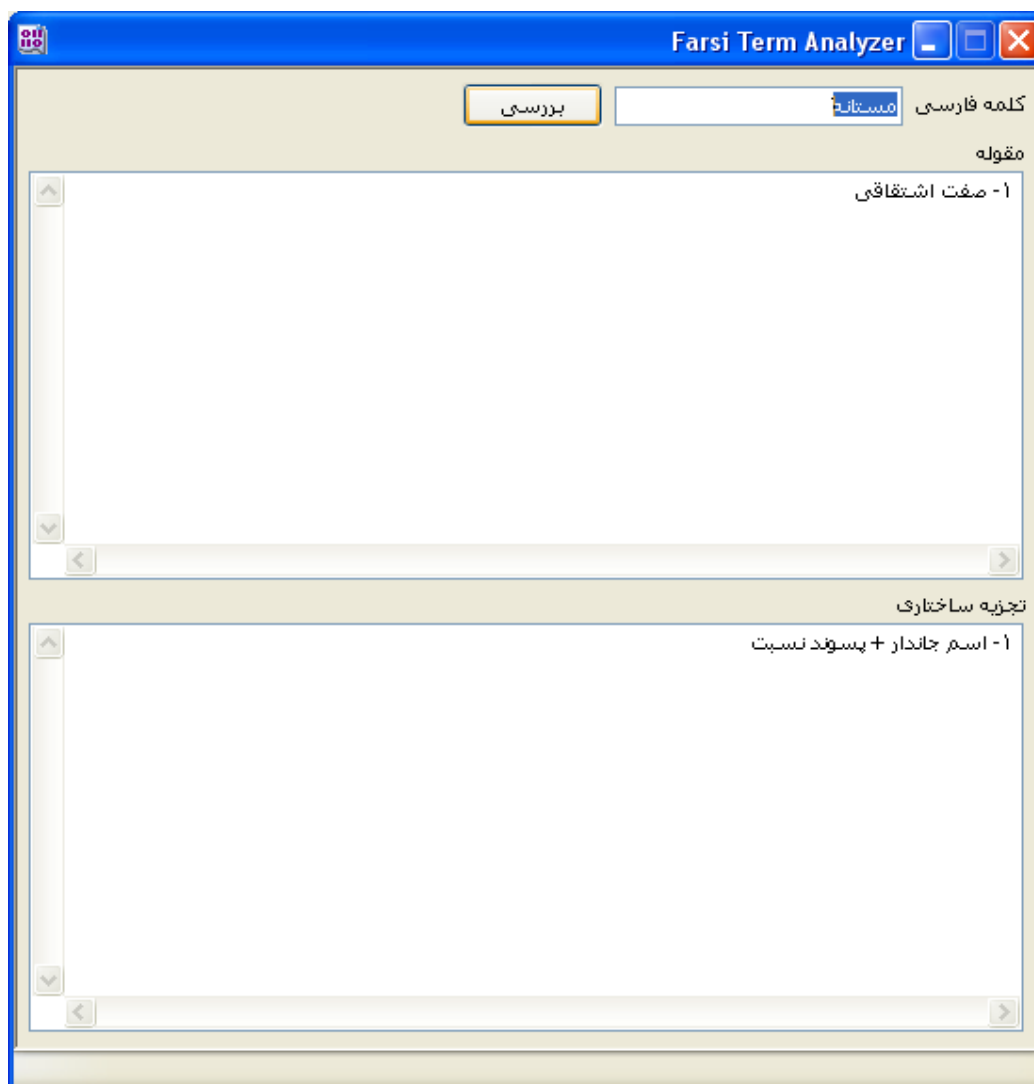












یادآوری: واژه شکن هر کلمه پیچیده را در دو خانه تحلیل می کند (خانه بالا و خانه پایین) در ردیف اول خانه بالا مقوله دستوری و در ردیف اول خانه پایین شکل تجزیه شده آن را می دهد. واژه شکن البته در اغلب داده ها به این اکتفا نمی کند و تحلیل آن را ارائه می دهد. در این حالت فقط ردیف اول خانه بالا و ردیف اول خانه پایین معتبرند و باقی

تحلیل‌ها مورد نظر نیستند. گاهی کلمه‌ای دو مقوله دستوری دارد که هر دو معتبرند بدیهی است در چنین حالتی اطلاعات هر دو ردیف مربوط به کلمه خانه بالا و پایین معتبرند. البته تحلیل‌های اضافی واژه شکن نیز بی فایده نیستند و قابلیت نرم افزار را ثابت می‌کند. تحلیل‌های اضافی با علامت ستاره مشخص شده‌اند.

برای هر کلمه مقوله دستوری و صورت تجزیه شده آن می‌آید. مثلاً خانه‌نشین اسم جاندار بن ساختی است یعنی از اسم و بن مضارع ساخته شده است که در بالا مقوله دستوری (اسم جاندار بن ساختی) و در پایین شکل تجزیه شده (اسم + بن مضارع) آمده است.

واژه شکن کلمه دسترنج را اسم جاندار بن ساختی می‌داند زیرا این ترکیب شامل اسم (دست) و بن مضارع (رنج) است. یعنی طبق قاعده و به قیاس با زورگیر و خانه‌نشین می‌توان دسترنج را اسم جاندار دانست. واژه شکن در مرحله بعد دسترنج را شامل قاعده اسم + اسم معنی می‌داند و آن را اسم مرکب می‌نامد.

صورت‌هایی که هم اسم و هم بن فعل‌اند با هر دو ویژگی در فهرست قواعد آمده‌اند ولی چون قاعده بن مضارع + بن مضارع معتبر نیست واژه شکن خواب‌نما را شامل اسم عمل + بن مضارع می‌داند و آن را اسم جاندار می‌گیرد.

خلاصه

واژه‌شکن کلمات پیچیده فارسی را تجزیه و مقوله آن را مشخص می‌کند. واژه‌شکن ساخت‌های بالقوه خوش‌ساخت را دستوری می‌داند و مقوله آن را به دست می‌دهد هرچند که ممکن است این ساخت‌ها از لحاظ معنایی غیر متعارف باشد. واژه‌شکن بدین ترتیب ساخت‌هایی نظیر چوب‌نشین و معلمستان را دستوری می‌داند و مقوله دستوری آن را مشخص و آن را به اجزاء سازنده‌اش تجزیه می‌کند.

قواعد دوری در واژه‌شکن عمل نمی‌کند. هر قاعده فقط یک بار عمل می‌کند و تکرار نمی‌شود. از ترکیب کتاب + پسوند تصغیر لفظی نظیر کتابک و از ترکیب اسم جاندار + پسوند مبالغه لفظ کبوتر باز ساخته می‌شود ولی کتابک که خود اسم و کبوتر باز که خود اسم جاندار است بار دیگر با پسوند تصغیر "ک" و پسوند مبالغه "باز" ترکیب نمی‌شوند تا ترکیبی دوری ساخته نشود.

۲-۳-۳-۱-۲- فعل شکن

فعل شکن برنامه‌ای رایانه‌ای است، که می‌تواند فعل‌های ساده را تشخیص داده و تجزیه کند. در این طرح فعل‌های مرکب و پیشوندی بررسی نشده‌اند. زیرا با حذف همکرد در فعل‌های مرکب و عنصر غیر فعلی در فعل‌های پیشوندی، آنچه باقی می‌ماند در واقع یک فعل ساده است. چگونگی بررسی گروه فعلی در بخش نحو قابل بررسی است.

قواعد ساخت صیغگان افعال

قاعده ۱، قاعده ساخت صفت مفعولی است که در ساخت ماضی نقلی و ماضی بعید و ماضی التزامی به کار می‌رود.

۱	بن ماضی + ه = صفت مفعولی
---	--------------------------

قواعد شماره ۲ تا ۸ و نیز ۱۱ قواعد کلی ساخت زمان‌های پر کاربرد زبان فارسی یعنی حال ساده، حال التزامی، ماضی نقلی، گذشته ساده، گذشته استمراری، ماضی بعید، ماضی التزامی و آینده را نشان می‌دهد. همان‌طور که ملاحظه می‌شود این قواعد کلی‌اند و مثلاً در قاعده ۲: می + بن مضارع + شناسه مضارع^۱ = حال ساده با جایگزینی هر یک از شناسه‌ها این صیغه برای تمامی شخص‌ها (اول، دوم، سوم) و شمارها (مفرد، جمع) صرف می‌شود.

از آنجایی که این قواعد برای کاربرد در نرم‌افزار نوشته شده است، نیازی به تعریف چگونگی اشتقاق بن ماضی و مضارع از یکدیگر نداریم زیرا فهرست این بن‌ها در پایگاه داده‌های نرم‌افزار وجود دارد که هریک از آن‌ها به صورت مستقل در فرآیند ساخت فعل

۱ نوشتن قواعد کلی با رویکرد کاربرد در ترجمه ماشینی باعث شد به سه دسته شناسه قائل شویم: شناسه‌های مضارع، شناسه‌های ماضی و شناسه‌های ماضی نقلی.

شرکت می کنند. نرم افزار مفروض قرار است قادر به تجزیه ساختار فعل باشد. که در این صورت هم نیازی به تعریف چگونگی اشتقاق بن ها از یکدیگر نیست.

۲	می + بن مضارع + شناسه مضارع = حال ساده
۳	ب + بن مضارع + شناسه مضارع = حال التزامی
۴	صفت مفعولی + شناسه نقلی = ماضی نقلی
۵	بن ماضی + شناسه ماضی = گذشته ساده
۶	می + بن ماضی + شناسه ماضی = گذشته استمراری
۷	صفت مفعولی + بود + شناسه ماضی = ماضی بعید
۸	صفت مفعولی + باش + شناسه مضارع = ماضی التزامی

صیغه امر تنها برای دوم شخص مفرد صرف می شود که به صورت قاعده ۹ نوشته شده است. ۱۰ در واقع قاعده ساخت امر منفی یا همان نهی است.

۹	ب + بن مضارع = امر دوم شخص مفرد
۱۰	ن + بن مضارع = نهی دوم شخص مفرد
۱۱	خواه + شناسه مضارع + بن ماضی = آینده

همه صیغه ها می توانند صورت منفی داشته باشند. در یک قاعده کلی می توان گفت با اضافه کردن پیشوند منفی ساز "ن" در ابتدای هر صیغه صورت منفی آن به دست می آید

(قاعده ۱۲). قواعد ۱-۱۲ تا ۷-۱۲ صورت تفصیلی قاعده ۱۲ اند. تنها ساخت حال التزامی

منفی از این قاعده مستثنی است که به صورت قاعده ۱۳ نوشته شد.

۱۲	۱-ن+ فعل = فعل منفی
۱-۱۲	ن + حال ساده = حال ساده منفی
۲-۱۲	ن + صفت مفعولی + شناسه نقلی = ماضی نقلی منفی
۳-۱۲	ن + بن ماضی + شناسه ماضی = گذشته ساده منفی
۴-۱۲	ن + می + بن ماضی + شناسه ماضی = گذشته استمراری منفی
۵-۱۲	ن + صفت مفعولی + بود + شناسه ماضی = ماضی بعید منفی
۶-۱۲	ن + صفت مفعولی + باش + شناسه مضارع = ماضی التزامی منفی
۷-۱۲	ن + خواه + شناسه مضارع + بن ماضی = آینده منفی
۱۳	۲-ن + بن مضارع + شناسه مضارع = حال التزامی منفی

برای اینکه بتوان قواعد را در برنامه نویسی کامپیوتری به کار برد، لازم است که با یک زبان میانجی آن‌ها را کدگذاری و سپس قواعد را با استفاده از آن سیستم کدگذاری بازنویسی کرد. در جدول‌های زیر این مراحل انجام شده است.

بن ماضی	S 1
---------	-----

بن مضارع	S 2
----------	-----

x در H_x ، G_x و N_x متغیری است که می‌تواند از ۱ تا ۶ متغیر باشد. و به ترتیبی که در جدول دیده می‌شود بر شخص و شمارهای مختلف دلالت کند. پس وقتی در قاعده‌ای می‌نویسیم H_x یعنی هر یک از صورت‌های آن (از H_1 تا H_6) می‌تواند به کار رود.

شناسه های مضارع H_x					
م	اول شخص مفرد	H_1	یم	اول شخص جمع	H_4
ی	دوم شخص مفرد	H_2	ید	دوم شخص جمع	H_5
د	سوم شخص مفرد	H_3	ند	سوم شخص جمع	H_6

شناسه های ماضی G_x					
م	اول شخص مفرد	G_1	یم	اول شخص جمع	G_4
ی	دوم شخص مفرد	G_2	ید	دوم شخص جمع	G_5
Ø	سوم شخص مفرد	G_3	ند	سوم شخص جمع	G_6

شناسه های ماضی نقلی N_X					
ام	اول شخص مفرد	N_1	ایم	اول شخص جمع	N_4
ای	دوم شخص مفرد	N_2	اید	دوم شخص جمع	N_5
است	سوم شخص مفرد	N_3	اند	سوم شخص جمع	N_6

متغیر X در شکل نهایی فعل همان شخص و شمار را نشان می‌دهد و H_X ، G_X و N_X تعیین کننده آن هستند. به عبارت دیگر X در V_{YX} ، با X در H_X ، G_X یا N_X سازنده آن برابر است. برای مثال در قاعده $V_{AX} = V_{A1} + S2 + H_X$ اگر $H_X = H_1$ باشد (یعنی شناسه اول شخص مفرد) آن گاه $V_{AX} = V_{A1}$ (یعنی فعل حال ساده اول شخص مفرد).

متغیر Y در V_{YX} ، بر زمان فعل دلالت دارد. همان‌طور که در جدول زیر پیداست Y می‌تواند A باشد و بر حال ساده دلالت کند، B باشد و بر ماضی نقلی دلالت کند، C باشد و بر گذشته ساده دلالت کند، D باشد و بر گذشته استمراری دلالت کند، E باشد و بر ماضی بعید دلالت کند، F باشد و بر ماضی التزامی دلالت کند، G باشد و بر آینده دلالت کند. در حال التزامی (V_{KX})، امر (V_{LX}) و نهی هم متغیر اول در واقع همان Y است، اما در جدولی جداگانه قرار گرفته‌اند زیرا ساخت منفی آن‌ها با بقیه زمان‌های ذکر شده تفاوت دارد (قاعده حال التزامی منفی N + بن مضارع + شناسه مضارع = حال التزامی)

منفی است و نهی صورت نفی امر است که تنها برای دوم شخص مفرد صرف می شود و

قاعده آن ن+ بن مضارع = نهی دوم شخص مفرد است و قاعده کلی ن+ فعل = فعل منفی

شامل آن نمی شود.

زمان V_{YX}			
V_{EX}	ماضی بعید	V_{AX}	حال ساده
V_{FX}	ماضی التزامی	V_{BX}	ماضی نقلی
V_{GX}	آینده	V_{CX}	گذشته ساده
		V_{DX}	گذشته استمراری

V_{KX}	حال التزامی
V_{LX}	امر
V_{JX}	نهی

P1	می
P2	ب
P3	ن
P4	ه

با در نظر داشتن آن چه پیش تر توضیح داده شد، فعل منفی لفظی است که بر کل صورت-

های منفی دلالت می کند. اگرچه حال التزامی منفی هم نوعی فعل منفی است، چون قاعده

ساخت آن با سایر ساخت‌های منفی متفاوت است لازم است با لفظی متفاوت نام گذاری شود.

$VN_Y X$	فعل منفی
$VN_K X$	حال التزامی منفی

A	صفت مفعولی
---	------------

F1	بود
F2	باش
F3	خواه

در جدول‌های زیر قواعد و صورت کدنویسی شده آن‌ها نشان داده شده‌است.

$S1 + P4 = A$	بن ماضی + ه = صفت مفعولی	۱
---------------	--------------------------	---

$P1 + S2 + H_X = V_A$	می + بن مضارع + شناسه مضارع = حال ساده	۲
$P2 + S2 + H_X = V_K X$	ب + بن مضارع + شناسه مضارع = حال التزامی	۳
۱-۳ استثناء: ب + ی + بن مضارعی که با آ آغاز می‌شود + شناسه مضارع = فعل حال التزامی		
$A + N_X = V_B X$	صفت مفعولی + شناسه نقلی = ماضی نقلی	۴
$S1 + G_X = V_C X$	بن ماضی + شناسه ماضی = گذشته ساده	۵

$P1 + S1 + G_X = V_{DX}$	می + بن ماضی + شناسه ماضی = گذشته استمراری	۶
$A + F1 + G_X = V_{EX}$	صفت مفعولی + بود + شناسه ماضی = ماضی بعید	۷
$A + F2 + H_X = V_{FX}$	صفت مفعولی + باش + شناسه مضارع = ماضی الترامی	۸

$P2 + S2 = V_{L2}$	ب + بن مضارع = امر دوم شخص مفرد	۹
$P3 + S2 = V_{J2}$	ن + بن مضارع = نهی دوم شخص مفرد	۱۰
$F3 + H_X + S1 = V_{GX}$	خواه + شناسه مضارع + بن ماضی = آینده	۱۱

$P3 + V_y x = V_{N_{YX}}$	ن + فعل = فعل منفی	۱۲
$P3 + V_a x = V_{N_{AX}}$	ن + حال ساده = حال ساده منفی	۱-۱۲
$P3 + A + N_X = V_{N_{BX}}$	ن + صفت مفعولی + شناسه نقلی = ماضی نقلی منفی	۲-۱۲
$P3 + S1 + G_X = V_{N_{CX}}$	ن + بن ماضی + شناسه ماضی = گذشته ساده منفی	۳-۱۲
$P3 + P1 + S1 = V_{N_{DX}}$	ن + می + بن ماضی + شناسه ماضی = گذشته استمراری منفی	۴-۱۲
$P3 + A + F1 + G_X = V_{N_{EX}}$	ن + صفت مفعولی + بود + شناسه ماضی = ماضی بعید منفی	۵-۱۲
$P3 + A + F2 + H_X = V_{N_{FX}}$	ن + صفت مفعولی + باش + شناسه مضارع = ماضی التزامی منفی	۶-۱۲

$P3 + F3 + Hx + S1 = VN_{GX}$	ن + خواه + شناسه مضارع + بن ماضی = آینده منفی	۷-۱۲
$P3 + S2 + Hx = VN_{KX}$	ن + بن مضارع + شناسه مضارع = حال التزامی منفی	۱۳

۲-۳-۳-۲-نحو

قواعد صوری دستور زبان مقصد یکی از ملزومات ترجمه ماشینی است. در ترجمه باید برای تمامی حوزه نحو زبان مقصد قواعدی صوری استخراج کرد. این قواعد که بسیاری از آنها مستقل اند و می توان آنها را بدون ارتباط با یکدیگر مشخص کرد، در نهایت به هم می پیوندند وقاعده ای منسجم برای نحو زبان مقصد می سازند.

سرفصل مواردی که باید در حوزه نحو مورد مطالعه دقیق قرار بگیرند عبارت اند از:

- انواع جمله
- ساختار درونی جمله
- گروه صفتی
- گروه قیدی
- گروه حرف اضافه ای
- گروه اسمی

- ضمائر

- همپایگی

- پیروسازی

هر یک از سرفصل‌های بالا که موضوع کلی یک طرح پژوهشی است، به قسمت‌های مختلفی تقسیم می‌شود، برای مثال:

- انواع جمله: در این قسمت باید به جملات خبری، پرسشی، امری و نیز نقل قول‌ها

پرداخته شود. هریک از این عنوان‌ها نیز، خود دارای عناوین فرعی است. برای

مثال در مورد جملات پرسشی، انواع جمله پرسشی، عناصر پرسشی‌ساز و ...

- ساختار درونی جمله: جمله ربطی، فعلی و انواع قید گونه‌ها

تمام موضوعات معرفی شده باید از دیدگاه کاربرد در ترجمه ماشینی مورد مطالعه

قرار بگیرد. یعنی خروجی هر طرح باید داده‌های مورد نیاز تجزیه گرنحوی نرم‌افزار ترجمه

را تامین کند. دقیقاً مانند واژه‌شکن و فعل‌شکن که خروجی آن‌ها برای کاربرد در تحلیل-

گر صرفی نرم‌افزار ترجمه طراحی شده است. در ادامه برای نشان دادن چگونگی مطالعه

مباحث نام‌برده، به معرفی پژوهشی درباره وابسته‌های پیشین گروه اسمی، که توسط

نگارنده انجام شده است، می‌پردازیم.

۲-۳-۳-۱- وابسته‌های پیشین گروه اسمی

گروه اسمی یکی از مهم‌ترین سازه‌های نحوی زبان فارسی و اغلب زبان‌های دنیا است که نقش مهمی در قواعد نحوی دارد. هر گروه اسمی یک هسته و یک مجموعه وابسته دارد. هسته گروه اسمی معمولاً اسم یا ضمیر است. وجود هسته اجباری و وابسته‌ها اختیاری است. اگر گروه اسمی فقط از یک کلمه ساخته شده باشد همان کلمه هسته گروه اسمی خواهد بود.

وابسته‌های پیش از هسته را وابسته‌های پیشین، و وابسته‌های پس از هسته را وابسته‌های پسین می‌نامند. وابسته‌های پیشین گروه اسمی به طبقات مشخصی نظیر صفت‌های اشاره و اعداد و صفت مبهم و صفت عالی و صفت‌های پرسشی و تعجبی و شاخص‌ها تقسیم می‌شوند. این طبقات هریک دارای اعضایی هستند. اعضای طبقات مذکور با قواعدی روی محور هم‌نشینی در کنار هم قرار می‌گیرند و قاعده‌های باهم‌آیی خاصی دارند. یعنی هم از لحاظ تقدم و تاخر و هم از نظر نوع وابسته‌ها تابع قاعده هستند. از طرفی برخی از این عناصر مانع الجمع‌اند.

وابسته‌های پیشین زبان فارسی، در منابع مختلف به طور اجمالی، به صورت زیر ارائه شده-

اند:

صادقی و ارژنگ (۱۳۵۶)	مشکوة الدینی (۱۳۷۳)	حق شناس و دیگران (۱۳۸۵)	سمائی (۱۳۸۱)
شاخص ها	وابسته وصفی	لقب ها (شاخص ها)	شاخص
اعداد اصلی (و ممیز)	وابسته عدد و ممیز	عددهای اصلی (و ممیز)	عدد (و ممیز) (شامل اصلی و)
عددهای ترتیبی و	وابسته عدد ترتیبی	عددهای ترتیبی و	ترتیبی)
صفت های عالی	وابسته صفت برترین	صفت های عالی	صفت عالی
صفت های اشاره	وابسته های اشاره	صفت های اشاره	صفت اشاره
صفت های پرسشی	وابسته های پرسشی	صفت های پرسشی	صفت پرسشی
صفت های تعجبی	وابسته های تعجبی	صفت های تعجبی	صفت تعجبی
صفت های مبهم	وابسته نامشخص	صفت های مبهم	صفت مبهم
یک نکره		یک نشانه نکره	یک نکره
		صفت ها بیانی	

در طرحی که به بررسی وابسته‌های پیشین در زبان فارسی، با نگاه کاربرد در ترجمه

ماشینی می‌پردازد، این وابسته‌ها به ترتیب جدول^۱ زیر دسته‌بندی شده‌اند:

^۱ در رابطه با جدول شماره ۱ توجه به چند نکته ضروری است:

۱- صفت‌های تعجبی بیشتر در سبک محاوره کاربرد دارد و در این طرح بررسی نشده‌اند.

۲- بما، بسی، فلان و بعضی دیگر از اعضای طبقات به دلایل سبکی حذف شده‌اند.

۳- "یک نکره" به دلایل سبکی و نیز این که باید با وابسته پسین "ی" بررسی شود، حذف شده‌است.

۴- شاخص به دلیل هم مقوله بودن با اسامی و لزوم اختصاص مولفه های معنایی انتزاعی به آن‌ها از فهرست قواعد حذف شده اند.

۵- دو نوع عدد ترتیبی وجود دارد که در طرح واژه شکن و عددشکن از آن‌ها به عنوان عددترتیبی نوع اول و عدد ترتیبی نوع دوم نام برده شده است. قاعده اعداد ترتیبی نوع اول و دوم در واژه شکن به ترتیب به صورت عدد اصلی + وند ترتیب^۱ (پنجم) و عدد اصلی + وند ترتیب^۲ (پنجمین) است. عدد ترتیبی نوع اول وابسته پسین و عدد ترتیبی نوع دوم وابسته پیشین است. در این طرح منظور از عدد ترتیبی، عدد ترتیبی نوع دوم است.

جدول شماره ۱:

صفت اشاره		صفت مبهم		صفت پرسشی		اعداد	ممیز	صفت عالی
بسیط	این	بسیط	هر	بسیط	چند	اعداد اصلی	عدد	صفت + ترین
	آن		هیچ		چندمین	اعداد ترتیبی		یگانه
	همین		همه		چه			تنها
	همان		بعضی		چگونه			
مرکب	چنین	مرکب	برخی	مرکب	چه نوع			
	چنان		چند		چقدر			
	یک		چندین		کدام			
	چنین							
	اینهمه	مرکب	بعضی از	مرکب	کدامیک از			
			برخی از					
			هریک از					
			هر کدام از					
			هیچیک از					
			هیچکدام از					

وابسته های پیشین برای قرار گرفتن در کنار هم محدودیت هایی دارند. یعنی هر وابسته ای ممکن است قبل و یا بعد از بعضی وابسته ها بیاید و در کنار بعضی دیگر قرار نگیرد. هم نشینی وابسته های پیشین ابتدا به صورت دو به دو بررسی شد که حاصل آن در جدول شماره ۲ ارائه شده است:

جدول شماره ۲:

صفت اشاره بسیط + اعداد اصلی	صفت اشاره + صفت پرسشی *	صفت اشاره + صفت مبهم *
صفت اشاره مرکب + اعداد اصلی *	(صفت اشاره بسیط + صفت پرسشی بسیط *) (صفت اشاره بسیط + صفت پرسشی مرکب *) (صفت اشاره مرکب + صفت پرسشی بسیط *) (صفت اشاره مرکب + صفت پرسشی مرکب *)	(صفت اشاره بسیط + صفت مبهم بسیط *) (صفت اشاره بسیط + صفت مبهم مرکب *) (صفت اشاره مرکب + صفت مبهم بسیط *) (صفت اشاره مرکب + صفت مبهم مرکب *)
صفت اشاره + اعداد ترتیبی نوع اول * (صفت اشاره بسیط + اعداد ترتیبی نوع اول *) (صفت اشاره مرکب + اعداد ترتیبی نوع اول *)	صفت پرسشی بسیط + صفت اشاره * (صفت پرسشی بسیط + صفت اشاره بسیط *) (صفت پرسشی بسیط + صفت اشاره مرکب *)	صفت اشاره بسیط + همه / چند
اعداد + صفت اشاره * (اعداد اصلی + صفت اشاره بسیط *) (اعداد اصلی + صفت اشاره مرکب *) (اعداد ترتیبی نوع اول + صفت اشاره بسیط *) (اعداد ترتیبی نوع اول + صفت اشاره مرکب *)	صفت پرسشی مرکب + صفت اشاره بسیط	صفت مبهم بسیط + صفت اشاره * (صفت مبهم بسیط + صفت اشاره بسیط *) (صفت مبهم بسیط + صفت اشاره مرکب *)

مرکب *		
	صفت پرسشی مرکب + صفت اشاره مرکب *	
		صفت مبهم مرکب + صفت اشاره بسیط
		صفت مبهم مرکب + صفت اشاره مرکب *
صفت اشاره بسیط + اعداد اصلی	صفت اشاره + صفت پرسشی * (صفت اشاره بسیط + صفت پرسشی بسیط *) (صفت اشاره بسیط + صفت پرسشی مرکب *) (صفت اشاره مرکب + صفت پرسشی بسیط *) (صفت اشاره مرکب + صفت پرسشی مرکب *)	صفت اشاره + صفت مبهم * (صفت اشاره بسیط + صفت مبهم بسیط *) (صفت اشاره بسیط + صفت مبهم مرکب *) (صفت اشاره مرکب + صفت مبهم بسیط *) (صفت اشاره مرکب + صفت مبهم مرکب *)
صفت اشاره مرکب + اعداد اصلی *		
صفت اشاره + اعداد ترتیبی نوع اول * (صفت اشاره بسیط + اعداد ترتیبی نوع اول *) (صفت اشاره مرکب + اعداد ترتیبی نوع اول *)	صفت پرسشی بسیط + صفت اشاره * (صفت پرسشی بسیط + صفت اشاره بسیط *) (صفت پرسشی بسیط + صفت اشاره مرکب *)	صفت اشاره بسیط + همه / چند
اعداد + صفت اشاره * (اعداد اصلی + صفت اشاره بسیط *)	صفت پرسشی مرکب + صفت اشاره بسیط	صفت مبهم بسیط + صفت اشاره * (صفت مبهم بسیط + صفت اشاره بسیط *)

(اعداد اصلی + صفت اشاره مرکب*) (اعداد ترتیبی نوع اول + صفت اشاره بسیط*) (اعداد ترتیبی نوع اول + صفت اشاره مرکب*)		(صفت مبهم بسیط + صفت اشاره مرکب*)
	صفت پرسشی مرکب + صفت اشاره مرکب*	
		صفت مبهم مرکب + صفت اشاره بسیط
		صفت مبهم مرکب + صفت اشاره مرکب*
صفت اشاره بسیط + اعداد اصلی	صفت اشاره + صفت پرسشی* (صفت اشاره بسیط + صفت پرسشی بسیط*) (صفت اشاره بسیط + صفت پرسشی مرکب*) (صفت اشاره مرکب + صفت پرسشی بسیط*) (صفت اشاره مرکب + صفت پرسشی مرکب*)	صفت اشاره + صفت مبهم* (صفت اشاره بسیط + صفت مبهم بسیط*) (صفت اشاره بسیط + صفت مبهم مرکب*) (صفت اشاره مرکب + صفت مبهم بسیط*) (صفت اشاره مرکب + صفت مبهم مرکب*)
صفت اشاره مرکب + اعداد اصلی*		
صفت اشاره + اعداد ترتیبی نوع اول* (صفت اشاره بسیط + اعداد ترتیبی نوع اول*)	صفت پرسشی بسیط + صفت اشاره* (صفت پرسشی بسیط + صفت اشاره بسیط*) (صفت پرسشی بسیط + صفت اشاره مرکب*)	صفت اشاره بسیط + همه/چند

(صفت اشاره مرکب + اعداد ترتیبی نوع اول (*)		
اعداد + صفت اشاره* (اعداد اصلی + صفت اشاره بسیط*) (اعداد اصلی + صفت اشاره مرکب*) (اعداد ترتیبی نوع اول + صفت اشاره بسیط*) (اعداد ترتیبی نوع اول + صفت اشاره مرکب*)	صفت پرسشی مرکب + صفت اشاره بسیط	صفت مبهم بسیط + صفت اشاره* (صفت مبهم بسیط + صفت اشاره بسیط*) (صفت مبهم بسیط + صفت اشاره مرکب*)
	صفت پرسشی مرکب + صفت اشاره مرکب*	
		صفت مبهم مرکب + صفت اشاره بسیط
		صفت مبهم مرکب + صفت اشاره مرکب*
صفت پرسشی بسیط + اعداد* (صفت پرسشی بسیط + اعداد اصلی*) (صفت پرسشی بسیط + اعداد ترتیبی*) کدام + اعداد	صفت مبهم + صفت پرسشی* (صفت مبهم بسیط + صفت پرسشی بسیط*) (صفت مبهم بسیط + صفت پرسشی مرکب*) (صفت مبهم مرکب + صفت پرسشی بسیط*) (صفت مبهم مرکب + صفت پرسشی مرکب*)	صفت اشاره بسیط + صفت عالی

(کدام + اعداد اصلی)		
(کدام + اعداد ترتیبی)		
صفت پرسشی مرکب + اعداد اصلی	صفت پرسشی + صفت مبهم* (صفت پرسشی بسیط + صفت مبهم بسیط*) (صفت پرسشی بسیط + صفت مبهم مرکب*) (صفت پرسشی مرکب + صفت مبهم بسیط*) (صفت پرسشی مرکب + صفت مبهم مرکب*)	صفت اشاره مرکب + صفت عالی*
صفت پرسشی مرکب + اعداد ترتیبی*		
	صفت مبهم + اعداد* (صفت مبهم بسیط + اعداد اصلی*) (صفت مبهم بسیط + اعداد ترتیبی نوع اول*) (صفت مبهم مرکب + اعداد اصلی*) (صفت مبهم مرکب + اعداد ترتیبی نوع اول*) (	صفت عالی + صفت اشاره* (صفت عالی + صفت اشاره بسیط*) (صفت عالی + صفت اشاره مرکب*)
صفت پرسشی بسیط + صفت عالی* صفت پرسشی مرکب + صفت عالی صفت پرسشی مرکب + یگانه / تنها*	صفت مبهم + صفت عالی* (صفت مبهم بسیط + صفت عالی*) (صفت مبهم مرکب + صفت عالی*)	

کدام + صفت عالی چند/چندمین + ممیز		
اعداد + صفت پرسشی* (اعداد اصلی + صفت پرسشی بسیط*) (اعداد ترتیبی نوع اول + صفت پرسشی مرکب*) (اعداد اصلی + صفت پرسشی بسیط*) (اعداد ترتیبی نوع اول + صفت پرسشی مرکب*)		
صفت عالی + صفت پرسشی* (صفت عالی + صفت پرسشی بسیط*) (صفت عالی + صفت پرسشی مرکب*)		
صفت پرسشی بسیط + اعداد* (صفت پرسشی بسیط + اعداد اصلی*) (صفت پرسشی بسیط + اعداد ترتیبی*) کدام + اعداد	صفت مبهم + صفت پرسشی* (صفت مبهم بسیط + صفت پرسشی بسیط*) (صفت مبهم بسیط + صفت پرسشی مرکب*) (صفت مبهم مرکب + صفت پرسشی بسیط*) (صفت مبهم مرکب + صفت پرسشی مرکب*)	صفت اشاره بسیط + صفت عالی

(کدام + اعداد اصلی)		
(کدام + اعداد ترتیبی)		
صفت پرسشی مرکب + اعداد اصلی	صفت پرسشی + صفت مبهم* (صفت پرسشی بسیط + صفت مبهم بسیط*) (صفت پرسشی بسیط + صفت مبهم مرکب*) (صفت پرسشی مرکب + صفت مبهم بسیط*) (صفت پرسشی مرکب + صفت مبهم مرکب*)	صفت اشاره مرکب + صفت عالی*
صفت پرسشی مرکب + اعداد ترتیبی*		
	صفت مبهم + اعداد* (صفت مبهم بسیط + اعداد اصلی*) (صفت مبهم بسیط + اعداد ترتیبی نوع اول*) (صفت مبهم مرکب + اعداد اصلی*) (صفت مبهم مرکب + اعداد ترتیبی نوع اول*) (	صفت عالی + صفت اشاره* (صفت عالی + صفت اشاره بسیط*) (صفت عالی + صفت اشاره مرکب*)
صفت پرسشی بسیط + صفت عالی* صفت پرسشی مرکب + صفت عالی صفت پرسشی مرکب + یگانه / تنها*	صفت مبهم + صفت عالی* (صفت مبهم بسیط + صفت عالی*) (صفت مبهم مرکب + صفت عالی*)	

کدام + صفت عالی چند/چندمین + ممیز		
اعداد + صفت پرسشی * (اعداد اصلی + صفت پرسشی بسیط *) (اعداد ترتیبی نوع اول + صفت پرسشی مرکب *) (اعداد اصلی + صفت پرسشی بسیط *) (اعداد ترتیبی نوع اول + صفت پرسشی مرکب *)		
صفت عالی + صفت پرسشی * (صفت عالی + صفت پرسشی بسیط *) (صفت عالی + صفت پرسشی مرکب *)		

سپس امکان هم نشینی همه وابسته های پیشین با یکدیگر مورد بررسی قرار گرفت .
جدول شماره ۳ شامل قواعد همه هم نشینی های مجاز است:

جدول شماره ۳:

صفت پرسشی مرکب + صفت اشاره بسیط + همه/چند	صفت پرسشی مرکب + صفت اشاره بسیط
صفت پرسشی مرکب + صفت اشاره بسیط + اعداد اصلی	صفت اشاره بسیط + اعداد اصلی
صفت پرسشی مرکب + صفت اشاره بسیط + اعداد اصلی + ممیز	صفت اشاره بسیط + همه/چند
صفت پرسشی مرکب + صفت اشاره بسیط + صفت عالی	صفت اشاره بسیط + صفت عالی
صفت پرسشی مرکب + اعداد اصلی + ممیز	صفت پرسشی مرکب + اعداد اصلی
	صفت پرسشی مرکب + صفت عالی
	صفت مبهم مرکب + صفت اشاره بسیط
کدام + اعداد + ممیز (کدام + اعداد اصلی + ممیز) (کدام + اعداد ترتیبی + ممیز)	اعداد + ممیز (اعداد اصلی + ممیز) (اعداد ترتیبی + ممیز)
	کدام + اعداد (کدام + اعداد اصلی) (کدام + اعداد ترتیبی)
	کدام + صفت عالی
	چند/چندمین + ممیز

برای استفاده این قواعد در نرم افزار نخست باید وابسته های پیشین کد گذاری شوند.
(جدول شماره ۵):

جدول شماره ۵:

A10	صفت اشاره
A10a	صفت اشاره بسیط
A10b	صفت اشاره مرکب
A11	صفت مبهم
A11a	صفت مبهم بسیط
A11b	صفت مبهم مرکب
A11a1	همه / چند
A12	صفت پرشی
A12a	صفت پرشی بسیط
A12b	صفت پرشی مرکب
A12a1	چند / چندمین
A12a2	کدام
A13	صفت عالی
A01S39	صفت + ترین
M	اعداد
M01	اعداد اصلی
A08 (M01S25)	اعداد ترتیبی
Q	ممیز

در رابطه با جدول شماره ۵ به نکات زیر توجه کنید:

- ۱- صفت عالی با همان کد گذاری مورد استفاده در واژه شکن کد گذاری شده است A13. قاعده ساخت صفت عالی در واژه شکن A01S39 است. A01 "صفت" و S39 پسوند " - ترین" است.
- ۲- اعداد هم با کد گذاری مورد استفاده در عدد شکن^۱ در این جا کد گذاری شده‌اند.

^۱ عدد شکن فارسی نرم افزاری است که قادر است گروه اعداد را پردازش کند و قواعد آن را معین کند و صورت دستوری و غیر دستوری آن را هم تمیز دهد. این نرم افزار نیز یکی از طرح های پژوهشی نگارنده است.

با کمک این زبان میانجی، قواعد برای استفاده در نرم افزار بازنویسی شد:

جدول شماره ۶:

صفت پرسشی مرکب + صفت اشاره بسیط + همه/چند A12b+ A10a+A11a1 صفت پرسشی مرکب + صفت اشاره بسیط + اعداد اصلی A12b+ A10a+M01 صفت پرسشی مرکب + صفت اشاره بسیط + اعداد اصلی + ممیز A12b+ A10a+M01+Q صفت پرسشی مرکب + صفت اشاره بسیط + صفت عالی A12b+ A10a+A13	صفت پرسشی مرکب + صفت اشاره بسیط A12b+ A10a صفت اشاره بسیط + اعداد اصلی A10a+M01 صفت اشاره بسیط + همه / چند A10a+ A11a1 صفت اشاره بسیط + صفت عالی A10a+A13
صفت پرسشی مرکب + اعداد اصلی + ممیز A12b+ M01+Q	صفت پرسشی مرکب + اعداد اصلی A12b+M01
	صفت پرسشی مرکب + صفت عالی A12b+ A13
	صفت مبهم مرکب + صفت اشاره بسیط A11b+ A10a
کدام + اعداد + ممیز A12a 2+M+Q (کدام + اعداد اصلی + ممیز) A12a 2+M01+Q)((کدام + اعداد ترتیبی + ممیز) (A12a 2+M01+A08)	اعداد + ممیز M+Q (اعداد اصلی + ممیز) (M01+Q) (اعداد ترتیبی + ممیز) (A08+Q)

	کدام + اعداد $A12a\ 2+M$ (کدام + اعداد اصلی) $(A12a\ 2+M01)$ (کدام + اعداد ترتیبی) $(A12a\ 2+A08)$
	کدام + صفت عالی $A12a\ 2+A13$
	چند / چندمین + ممیز $A12a1+Q$

جدول شماره ۷:

$A12b+ A10a$
$A10a+M01$
$A10a+ A11a1$
$A10a+A13$
$A12b+M01$
$A12b+ A13$
$A11b+ A10a$
$M+Q$ $(M01+Q)$ $(A08+Q)$
$A12a\ 2+M$ $(A12a\ 2+M01)$ $(A12a\ 2+A08)$
$A12a\ 2+A13$
$A12a1+Q$

A12b+ A10a+A11a1
A12b+ A10a+M01
A12b+ A10a+M01+Q
A12b+ A10a+A13
A12b+ M01+Q
A12a 2+M+Q
(A12a 2+M01+Q)
(A12a 2+M01+A08)

جمع‌بندی

فصل اول به مرور کلی سابقه پژوهش‌های انجام‌شده در زمینه مترجم ماشینی و معرفی نرم-افزارها اختصاص داده شد. در این بخش ابتدا به تعریف زبان‌شناسی رایانه‌ای و ترجمه ماشینی به عنوان یکی از فعالیت‌های این حوزه پرداخته شد.

در این فصل تاریخ ترجمه ماشینی بعد از جنگ جهانی دوم تا امروز از نظر گذراننده شد و منابع، روش‌ها، نظریه‌ها و فناوری‌های موجود و پیشرفت‌های صورت گرفته در این زمینه معرفی و نقد شد.

در فصل دوم به طرح‌های مرتبط با تحقق مترجم ماشینی پرداخته شد. در این فصل به تفصیل توضیح داده شد که برای داشتن یک مترجم ماشینی، زبان مبدا و مقصد باید به طور دقیق بررسی و قواعد آن استخراج شود. متناسب با موضوع طرح حاضر ("شناسایی طرح‌های مرتبط با تحقق مترجم ماشینی در زبان فارسی")، به بررسی زبان‌فارسی از دیدگاه ترجمه ماشینی پرداختیم. برای نشان‌دادن چگونگی استخراج قواعد، چند طرح که پیش‌تر توسط نگارنده انجام شده بود، معرفی شد. شایان ذکر است که با همکاری متخصصان رایانه، برای دو طرح از طرح‌های نام‌برده، نرم‌افزار تهیه شده است.

به طور کلی طرح‌های مرتبط با تحقق مترجم ماشینی در زبان فارسی عبارتند از:

- تدوین واژگان (با مشخصات تصریح‌شده در متن)

- قواعد ساخت انواع جمله پرسشی
- قواعد ساخت انواع جمله امری
- قواعد ساخت انواع جمله خبری
- قواعد ساخت جملات همپایه
- قواعد ساخت گروه اسمی
- قواعد ساخت گروه صفتی
- قواعد ساخت گروه قیدی
- قواعد ساخت گروه حرف اضافه‌ای
- قواعد مرجع‌گزینی
- قواعد پیروسازی با بند اسمی
- قواعد پیروسازی با بند صفتی
- قواعد پیروسازی با بند قیدی
- قواعد تطابق زمان‌ها

واژه‌نامه فارسی به انگلیسی

phonology	آواشناسی
lexical ambiguity	ابهام ساختاری
lexical ambiguity	ابهام واژگانی
word-sense disambiguation	ابهام‌زدایی از معنای واژه
derivation	اشتقاق
structural transfer	انتقال ساختاری
lexical transfer	انتقال واژه‌ها
context	بافت
heuristic	بینشی
ending	پایانه کلمات
base	پایه
bottom-up	پایین به بالا
post-editing	پس ویرایش
rules post-processed by statistics	پس‌ویرایش آماری قواعد
pre-editing	پیش‌ویرایش
corpus	پیکره
input analysis	تحلیل ورودی
word Order	ترتیب کلمات
interactive translation	ترجمه تعاملی
computer-aided translation	ترجمه رایانه‌ای
transfer-based machine translation	ترجمه ماشینی انتقال بنیان

hybrid machine translation	ترجمه ماشینی پیوندی
shallow-transfer machine translation	ترجمه ماشینی سطحی
dictionary-based machine translation	ترجمه ماشینی فرهنگ بنیان
example-based machine translation	ترجمه ماشینی مثال بنیان
interlingua machine translation	ترجمه ماشینی میان‌زبانی
compounding	ترکیب
inflectional	تصریفی
interactive	تعاملی
tokenization	واحدسازی
distribution	توزیع
location	جایگاه
substitution	جایگزین‌سازی
subcategorization frame	چارچوب زیرمقوله‌ای
polysemy	چند معنایی
linear	خطی
well-formed	خوش ساخت
universal encyclopedia	دائرةالمعارف جهانی
linguistics knowledge	دانش زبانی
synset	دایره معنایی
text understanding	درک متن
lexical Functional Grammar	دستور واژگانی نقشمند

recursive	دوری
controlled language	زبان مهارشده
light verb construction	ساختار فعل های سبک
constituent	سازه
parsing	سازه یابی
shallow	سطحی
motphology	صرف
citation form	صورت اسنادی
subworld	عالم فرعی
punctuation	علائم سجاوندی
deep	عمیق
extended dictionary	فرهنگ لغت گسترده
Light verb	فعل سبک
open source	متن باز
scenario	فیلمنامه
rule-based	قاعده بنیان
domain	قلمرو
subcategorization rules	قواعد زیرمقوله ای
analogy	قیاس
Key words	کلیدواژه ها
phrase	گروه
unrestricted	محدود نشده
boundary	مرز

semantics	معنی شناسی
Iterative	مکرر
formulaic	منضبط
intermediary	میانجی
syntax	نحو
orthographic	نگارشی
scripts	نمایشنامه
conceptual dependency	وابستگی مفهومی
lexicon	واژگان
lexico - syntactic	واژگانی - نحوی
gloss	واژه نامه
statistics guided by rules	ویرایش آمار با قواعد
homonymy	هم معنایی
homophone	هماوا

واژه‌نامه انگلیسی به فارسی

analogy	قیاس
base	پایه
bottom-up	پایین به بالا
boundary	مرز
citation form	صورت اسنادی
compounding	ترکیب
computer-aided translation	ترجمه رایانه‌ای
conceptual dependency	وابستگی مفهومی
constituent	سازه
context	بافت
controlled language	زبان مهارشده
corpus	پیکره
deep	عمیق
derivation	اشتقاق
dictionary-based machine translation	ترجمه ماشینی فرهنگ بنیان
distribution	توزیع
domain	قلمرو
ending	پایانه کلمات
example-based machine translation	ترجمه ماشینی مثال بنیان
extended dictionary	فرهنگ لغت گسترده
formulaic	منضبط

gloss	واژه‌نامه
heuristic	بینشی
homonymy	هم معنایی
homophone	هماوا
hybrid machine translation	ترجمه ماشینی پیوندی
inflectional	تصریفی
input analysis	تحلیل ورودی
interactive	تعاملی
interactive translation	ترجمه تعاملی
interlingua machine translation	ترجمه ماشینی میان‌زبانی
intermediary	میانجی
Iterative	مکرر
Key words	کلیدواژه‌ها
lexical ambiguity	ابهام ساختاری
lexical ambiguity	ابهام واژگانی
lexical Functional Grammar	دستور واژگانی نقشمند
lexical transfer	انتقال واژه‌ها
lexicon	واژگان
lexico - syntactic	واژگانی - نحوی
Light verb	فعل سبک
light verb construction	ساختار فعل‌های سبک
linear	خطی

linguistics knowledge	دانش زبانی
location	جایگاه
morphology	صرف
open source	متن باز
orthographic	نگارشی
parsing	سازه یابی
phonology	آواشناسی
phrase	گروه
polysemy	چند معنایی
post-editing	پس ویرایش
pre-editing	پیش ویرایش
punctuation	علائم سجاوندی
recursive	دوری
rule-based	قاعده بنیان
rules post-processed by statistics	پس ویرایش آماری قواعد
scenario	فیلمنامه
scripts	نمایشنامه
semantics	معنی شناسی
shallow	سطحی
shallow-transfer machine translation	ترجمه ماشینی سطحی
statistics guided by rules	ویرایش آمار با قواعد
structural transfer	انتقال ساختاری

subcategorization frame	چارچوب زیرمقوله ای
subcategorization rules	قواعد زیرمقوله ای
substitution	جایگزین سازی
subworld	عالم فرعی
synset	دایره معنایی
syntax	نحو
text understanding	درک متن
tokenization	واحدسازی
transfer-based machine translation	ترجمه ماشینی انتقال بنیان
universal encyclopedia	دائرةالمعارف جهانی
unrestricted	محدود نشده
well-formed	خوش ساخت
word Order	ترتیب کلمات
word-sense disambiguation	ابهام زدایی از معنای واژه

فهرست منابع

- بیجن خان، م و مرادزاده، ش.، (۱۳۸۳)، "پیکره متنی زبان فارسی"، مجموعه سخنرانی‌های اولین کارگاه پژوهشی زبان فارسی و رایانه، تهران.
- ثانی، غ و سلیمانی، ع، (۱۳۷۸)، "ترجمه نحوی انگلیسی به فارسی به روش تجزیه چارت"، زبان فارسی و رایانه، تهران.
- حق شناس، ع، (۱۳۸۵). دستور زبان فارسی، وزارت آموزش و پرورش.
- داودآبادی، م. و پالهنک، م.، (۱۳۸۴)، "دانا: عاملی با قابلیت درک و نوشتار فارسی"، یازدهمین کنفرانس بین‌المللی انجمن کامپیوتر ایران.
- سمائی، م.، (۱۳۷۷)، "واژگان در دستور سنج: انگاره نظری"، رساله دکتری، دانشکده ادبیات و علوم انسانی، دانشگاه تهران.
- ، (۱۳۸۱)، پردازش گروه اسمی، فصلنامه اطلاع رسانی، دوره ۱۸، شماره ۱ و ۲.
- ، (۱۳۸۴). گروه اعداد در ترجمه ماشینی. فصلنامه علوم و فناوری اطلاعات، ۲۱ (۱)، ۳۳-۲۱.
- سمیعی، ا. (مترجم)، (۱۳۶۲)، **ساخت های نحوی**، نوشته نعام چامسکی، انتشارات خوارزمی.
- شمس فرد، م.، (۱۳۷۴)، "درک متن فارسی"، پایان‌نامه کارشناسی ارشد، دانشکده مهندسی کامپیوتر، دانشگاه صنعتی شریف.
- شمس فرد، م.، (۱۳۸۵)، "پردازش متون فارسی: دستاوردهای گذشته، چالش‌های پیش‌رو"، مجموعه مقالات دومین کارگاه پژوهشی زبان فارسی و رایانه.
- صادقی، ع، (۱۳۵۶). دستور زبان فارسی (سال دوم فرهنگ و ادب)، انتشارات آموزش و پرورش.
- صامتی، ح و بیجن خان، م.، (۱۳۸۹) **زبان فارسی و رایانه**، سمت.
- فیلی، ه.، (۱۳۸۳)، "ترجمه ماشینی انگلیسی به فارسی بر مبنای نظریه احتمالات و گرامر لغوی"، مجموعه سخنرانی‌های اولین کارگاه پژوهشی زبان فارسی و رایانه، تهران.
- فیلی، ه و ثانی، غ، (۱۳۸۳)، "ترجمه ماشینی انگلیسی به فارسی بر مبنای نظریه احتمالات و دستور کلمه"، زبان فارسی و رایانه.
- مشکوٰۃ‌الدینی، م.، (۱۳۷۳). دستور زبان فارسی (بر پایه نظریه گشتاری)، انتشارات دانشگاه فردوسی.

ناصح، م.، (۱۳۸۵)، "فهرست پایان‌نامه‌های ارائه شده در زمینه زبان فارسی و رایانه، ۱۳۵۵-۱۳۸۴"، مجموعه مقالات دومین کارگاه پژوهشی زبان فارسی و رایانه.
 نوربالا، م.، (۱۳۸۰)، "طراحی یک واژگان جامع برای پردازشگر زبان فارسی"، پایان‌نامه کارشناسی ارشد، دانشگاه علم و صنعت.

- Amtrup, J., Mansouri Rad, H., Megerdooonian, K., Zajac. (2000). *Persian – English Machine Translation: An Overview of the Shiraz Project*, Memoranda in Computer and Cognitive Science.
 between English and Persian Texts, Third International Workshop on Computational Approaches to Arabic-Script based languages (CAASL3), Ottawa, Canada.
- Booth, K. H. V. (1967). Machine aided translation with a post-editor. In A. D. Booth (Ed.), *Machine translation* (pp. 53-76). Amsterdam: North-Holland Publishing Company.
- Crystal, D, 1992, an encyclopedic dictionary of language and languages, Blackwell.
- Delvaney, Emile (1972). *La machine a traduire*. Paris: PUF.
- Dreyfus, H. L. & Dreyfus, S. E. (1989). *Mind over machine*. Oxford , Uk: B. Blackwell.
- Fellbaum, C. 1998. *WordNet An Electronic Lexical Database*, MIT Press.
- Banerjee, S. and T. Pederson 2002. *An Adapted Lesk Algorithm for Word Sense Disambiguation Using WordNet*. Lecture Notes in Computer Science, Berlin/Heidelberg, Springer.
- Harabagiu, S. M., G. A. Miller, et al. 1999. *Wordnet 2 – a morphologically and semantically enhanced resource*, University of Maryland, SIGLEX-99.
- Katz, J., & Fodor, J. (1963). The structure of a semantic theory. *Language*, 39, 170 – 210.
- Megredoomian. M., 2000, *Persian Computational Morphology: AUnification-Based Approach*, Memoranda in Computer and Cognitive Science. New Mexico State University.
- Katz, J., & Postal, P. M. (1964). *An integrated theory of linguistic descriptions*. Cambridge: MIT press.
- Kiani. S. 2009, *Persian text tokenization and chunking*, 14th international CSI Computer Conference. Tehran. Iran.

- Lesk, M. (1986). *Automatic sense disambiguation using machine readable dictionaries: How to tell a pine cone from an ice cream cone*. SIGDOC 86:the 5th annual international conference on Systems documentation, New York, USA, ACM Press.
- LISA (Library and Information Science Abstracts). (1997).
- Mahmoudi, S. M. (1994). *Contribution au traitement automatique de la langue persan*. These de doctorat, Universite de Lyon II .
- Motazedi, Y., Saedi, C., Shamsfard, M., (2009). "Automatic Translation
- Rosseta, M.T. (1994). *Rosseta Compositional Translation*, Cluwer Academic Publishers.
- Rouhizadeh. M., M. Shamsfard and M. A. Yarmohamadi, 2008, Building a Wordnet for Persian Verbs, 4th Global WordNet conference (GWC 2008), Szeged, Hungary.
- Shamsfard. M., H. Sadat Jafari and M. Ilbeigy, 2010, *A Set of Fundamental Tools for Persian Text Processing*. Seventh conference on International Language Resources and Evaluation (LREC'10), Valletta, Malta.
- http://en.wikipedia.org/wiki/Machine_translation (accessed at 02.03.2011)

ⁱ Nagao, M. 1981. A Framework of a Mechanical Translation between Japanese and English by Analogy Principle, in Artificial and Human Intelligence, A. Elithorn and R. Banerji (eds.) North- Holland, pp. 173-180, 1984.

ⁱⁱ "the Association for Computational Linguistics - 2003 ACL Lifetime Achievement Award". Association for Computational Linguistics.
http://www.aclweb.org/index.php?option=com_content&task=view&id=36&Itemid=30. Retrieved 2010-03-10.

ⁱⁱⁱ Boretz, Adam, "AppTek Launches Hybrid Machine Translation Software" SpeechTechMag.com (posted 2 MAR 2009)

^{iv} Milestones in machine translation - No.6: Bar-Hillel and the nonfeasibility of FAHQT by John Hutchins

^v Bar-Hillel (1960), "Automatic Translation of Languages". Available online at <http://www.mt-archive.info/Bar-Hillel-1960.pdf>

^{vi} Melby, Alan. The Possibility of Language (Amsterdam:Benjamins, 1995, 27-41)

^{vii} Wooten, Adam. "A Simple Model Outlining Translation Technology" T&I Business (February 14, 2006)

^{viii} Appendix III of 'The present status of automatic translation of languages', Advances in Computers, vol.1 (1960), p.158-163. Reprinted in Y.Bar-Hillel: Language and information (Reading, Mass.: Addison-Wesley, 1964), p.174-179.

^{ix} Human quality machine translation solution by Ta with you