

 وزارت آموزش عالی و پژوهش	 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	

عنوان زیرپروژه:

بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی

 وزارت فرهنگ و آموزش عالی	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

فهرست مطالب

شماره صفحه

عنوان

1	مقدمه.....
3	فصل اول: ترجمه ماشینی.....
5	1-1- مقدمه ای بر ترجمه ماشینی.....
5	1-2- تاریخچه مختصر ترجمه ماشینی.....
10	1-3- انواع سیستم های ترجمه ماشینی.....
10	1-3-1- سیستم های مستقیم.....
12	1-3-2- سیستم های غیر مستقیم.....
12	1-3-2-1- روش انتقالی.....
14	1-3-2-1-1- انتقال واژگانی.....
15	1-3-2-1-1- انتقال ساختاری.....
15	1-3-2-2- روش زبان میانجی.....
16	1-3-3- روش های مبتنی بر قانون.....
17	1-3-4- روشهای مبتنی بر دانش.....

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	

- 20 1-3-5-1 ماشین های ترجمه مبتنی بر پیکره (زبانی)
- 21 1-3-5-1-1 روشهای مبتنی بر نمونه
- 23 1-3-5-2-1 ترجمه ماشینی آماری
- 25 1-3-5-2-1-1 پیکره متناظر واژه ای
- 26 1-3-5-2-2-1 پیکره متناظر جمله ای
- 26 1-3-5-2-2-2 مراحل ساخت یک پیکره در ماشین ترجمه آماری
- 27 الف - پیش پردازش و علامت گذاری متون
- 27 ب - هم ترازی
- 29 1- ب - اندازه گیری رخداد کلمات و عبارات
- 30 2- ب - یک استراتژی در هم ترازی عبارات
- 31 3- ب - طبقه بندی و انتخاب عبارات
- 32 4- ب - هم ترازی در سطح عبارات
- 32 5- ب - هم ترازی در سطح کلمات
- 33 6- ب - پس پردازش
- 34 1-3-6-1 ترجمه ماشینی مبتنی بر اصل
- 34 1-3-7-1 روشهای ترکیبی
- 35 1-4-1 مکانیزم، مؤلفه ها و اجزای سازنده سیستم های ترجمه ماشینی

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

- 36 1-4-1- واژگان
- 37 1-4-2- نحو و ساختار
- 38 1-4-3- معناشناسی
- 39 1-5- ابهام زدایی معنی واژه
- 40 1-5-1- روش های مبتنی بر دانش یا مبتنی بر قاعده
- 40 1-5-2- استفاده از روشهای آماری در ابهام زدایی معنایی
- 41 1-5-2-1- استفاده از گنجواژه
- 42 1-5-2-1-1- محدود سازی طبقه
- 42 1-5-2-1-2- جمع بارها
- 43 1-5-2-1-3- انتخاب محتمل ترین طبقه
- 43 1-5-3- استفاده از موصوف برای دستیابی به معنای صفت
- 45 1-5-4- ترکیب آماری شاخص های چند گانه
- 45 1-5-4-1- مدل های تجزیه شدنی
- 46 1-5-4-2- ابهام زدایی مبتنی بر نمونه
- 47 1-5-5- استفاده از پیکره یک زبانه زبان مقصد
- 47 1-5-5-1- روش داگان و ایتای
- 48 1-5-5-2- روش هیونگ لی و جنگ سی پارک و چنگ کیم

 وزارت آموزش عالی و تحقیقات علمی	 ثروای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

48	الف- آشکار سازی معنی کلمه مبدأ و نگاشت کلمه.....
49	ب-نگاشت کلمه زبان مبدأ به مجموعه ای از کلمات زبان مقصد.....
50	ج- انتخاب واژه با استفاده از اطلاعات آماری زبان مقصد.....
52	3-5-1- روش موسوی و دلاور خلفی.....
54	فصل دوم: مقدمه ای بر کشف دانش.....
54	1-2- مقدمه ای بر کشف دانش
54	2-2- بررسی موضوع و مسائل پیرامون آن.....
54	3-2- تعاریف اصلی.....
54	1-2-3- طبقه بندی چند هدفه و چند طبقه بودن.....
54	2-2-3- کشف دانش در پایگاه داده ها.....
54	3-2-3- داده کاوی.....
55	4-2-3- عامل.....
55	4-2- هدف پروژه.....
55	5-2- مروری بر کشف دانش در پایگاه داده ها.....
55	1-2-5- دلایل توجه.....
57	2-2-5- فناوری های کلیدی مورد استفاده.....

 جمهوری اسلامی ایران	عنوان پروژه:		 شرایع ملی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

فصل سوم: کشف دانش بوسیله الگوریتم بهینه سازی

58

58 3-1- کشف دانش بوسیله الگوریتم بهینه سازی

58 3-2- فعالیت طبقه بندی

59 3-3- مزیت‌های الگوریتم بهینه سازی

59 3-3-1- الگوریتم بهینه سازی

59 3-3-2- حل مسئله

59 3-4- الگوریتم بهینه سازی

63 3-5- استفاده از الگوریتم در حل مسئله

64 3-6- الگوریتم و جزئیات یادگیری مسئله

65 3-7- سیستم یادگیری

67 فصل چهارم: کشف دانش چند طبقه مبتنی بر بهینه سازی

67 4-1- کشف دانش چند طبقه مبتنی بر بهینه سازی

67 4-2- دانش چند طبقه و نحوه نمایش آن

68 4-3- تحلیل هدف از نقطه نظر الگوریتم بهینه سازی

70 4-3-1- گام‌های دیگر طراحی

 جمهوری اسلامی ایران	عنوان پروژه:			 شرایع ملی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه:			
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

71	4-4- سلسله مراتب در الگوریتم
71	4-4-1- پی آمد ناشی از تخصصی شدن کار الگوریتم ها
72	4-4-2- تغییر فضای مسئله
73	4-5- بررسی چگونگی انتخاب عبارات
74	4-6- نحوه کار الگوریتم چرخ رولت در انتخاب عبارت ها
74	4-7- الگوریتم کشف دانش چند طبقه مبتنی بر ACO
80	4-7-1- ساخت قانون
81	4-7-2- هرس قانون
81	4-7-3- بروز رسانی فرومون
82	4-7-4- ساخت قوانین طبقه بندی
83	4-7-5- تابع هیوریستیک
85	4-7-6- هرس قوانین طبقه بندی
86	4-7-7- قوانین انتقال و بروز رسانی فرومون
89	فصل پنجم: ابهام زدایی واژه
89	5-1- ابهام زدایی واژه با استفاده
94	5-2- نتایج و بررسی کارایی الگوریتم ابهام زدایی از معنی واژه

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	

94	5-3- نحوه طبقه بندی بوسیله قوانین کشف شده
94	5-4- مجموعه آموزشی
95	5-5- تنظیم پارامترهای الگوریتم
96	5-6- نتایج کسب شده
97	1-1-5-6- بررسی نتایج کسب شده
99	2-1-5-6- محاسن بکارگیری این روش در مقایسه با روشهای موجود
100	فصل ششم - نتیجه گیری
103	منابع و مراجع
107	پیوست

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیر پروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن‌فارس — 2 — پ	

مقدمه

پیشرفت علوم و فن آوری و پیرو آن ، راحتتر شدن زندگی امروزی ، پیوندی ذاتی و ناگسستنی با دانش بشری دارد و در این میان بشر از همان ابتدای خلقت سعی بر این داشته تا با الهام از طبیعت پیرامون خود برای مشکلات و موانعی که بر سر راه داشته راه چاره ای بیاندیشد. از طرفی حجم زیاد داده هایی که در زمینه های مختلف علوم بدست آمده اند که با رشدی تصاعدی و روز افزون ، در جریانی یک طرفه به سمت ما می آیند ، که این خود باعث بحرانی شدن مسئله دسترسی به دانش مورد نیاز گردیده است. این پدیده نیز به تبعه خود باعث بحرانی شدن فعالیتهایی مانند تصمیم گیری¹ و یا طرح ریزی² در حوزه های مرتبط با دانش می شود. (۱)

با این وجود، فناوری ابزارهای لازم برای مقابله با این مشکل را در اختیار ما نهاده است. تکنیکهای سخت افزاری ، نرم افزاری و هوش مصنوعی به عنوان راههای مقابله با آن می توانند در نظر گرفته شوند . که تلفیق این تکنیکها با هم نیز خود باعث شکل گیری راههای جدید شده است. از این میان جستجو و کشف دانش در پایگاه داده ها³ با مجموعه ای از روشهایی که امکان غربال کردن داده ها و بدست آوردن دانش را در اختیار ما می گذارد ، قدرت و انعطاف مناسبی را از خود نشان داده است.

از میان روشهای متعددی که برای کشف دانش در پایگاه داده ها وجود دارند ، روشی خاص مبتنی بر الگوریتم بهینه سازی ها⁴ (2)(3)(4) مطرح گردیده است .

هدف از پروژه ای که این پروژه گزارش تحلیلی آن می باشد ، استفاده از کشف دانش چند طبقه با بکارگیری الگوریتم بهینه سازی در ابهام زدایی از معنی واژه است. این الگوریتم توانایی زیادی در حل مسائل ترکیبی⁵ نشان داده است (3) لزوم انجام این پروژه مشخص می شود.

فصل اول بطور خلاصه به ترجمه ماشینی می پردازد و موضوع اصلی مورد تحقیق را مورد بررسی قرار میدهد. مقدمه ای بر الگوریتم بهینه سازی در فصل دوم مرور می گردد . کشف دانش بوسیله الگوریتم بهینه سازی و الگوریتم های بهینه سازی مبتنی بر آن و و مزایای این الگوریتم در فصل سوم مورد مطالعه قرار می گیرد. در فصل

¹ Decision making

² Plannig

³ Knowledge Discovery in Data Base(KDD)

⁴ Ant Colony Optimization(ACO)

⁵ Combinatorial

	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	کد زیرپروژه: پیک‌متن فارس — 2 — پ	ویرایش: 1/0	تاریخ: 1388/04/7

چهارم کشف دانش چند طبقه مبتنی بر بهینه سازی به عنوان یک جزء از روشهای جستجوی دانش بررسی می گردد. و ابهام زدایی معنی واژه و نتایج پیاده سازی این الگوریتم بر روی یک پیکره فارسی در فصل پنجم مطالعه می شود.

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

فصل اول: ترجمه ماشینی

1-1- مقدمه ای بر ترجمه ماشینی

مکانیزه کردن ترجمه یکی از نخستین رؤیاهای بشریت بوده است. در قرن بیستم این رؤیا به شکل برنامه های کامپیوتری که قادر به ترجمه حجم زیادی از متون از یک زبان به زبان دیگر بودند به حقیقت پیوست. اما چون همیشه، واقعیت کامل نیست. ماشینهای ترجمه ای وجود ندارند که با فشار چند تکمه هر متن به هر زبانی را دریافت و ترجمه کاملی به هر زبان دیگر بدون دخالت یا کمک انسان تولید کنند. این کمال مطلوبی است برای آینده ای نزدیک که اگر چه اساساً قابل دسترسی است، رسیدن به آن چندان قطعی نیست. آنچه که تاکنون به آن رسیده ایم ساخت برنامه هایی است که می توانند ترجمه های خامی از متون در موضوعات نسبتاً تعریف شده ای تولید کنند که بعداً مورد تجدید نظر قرار می گیرند تا متون ترجمه شده ای با کیفیت خوب و مناسب از نظر اقتصادی ارائه دهند یا اینکه به همان حالت اصلاح نشده به وسیله متخصصان آن موضوع صرفاً برای گرفتن اطلاعات خوانده و فهمیده شوند. در بعضی موارد با کنترلهای مناسب بر روی زبان متون ورودی، ترجمه های خودکاری با کیفیت بهتر که نیاز کمی به اصلاح دارند یا اصلاً چنین نیازی ندارند می توان تولید کرد.

اینها پیشرفتهایی هستند در آنچه که معمولاً ترجمه ماشینی نامیده می شوند، اما این مفهوم اغلب مبهم بوده و اشتباه تعبیر می شود. دریافت عمومی از ترجمه ماشینی به وسیله دو نقطه نظر افراطی تحریف شده است. از یک طرف کسانی هستند که متقاعد نشده اند که تجزیه زبان با مشکل همراه است چرا که حتی کودکان قادرند که براحتی زبان را فراگیرند. اینها کسانی هستند که معتقدند هر کس که یک زبان خارجی را بداند باید بتواند براحتی آن را ترجمه کند. بنابراین آنها قادر به فهم مشکلات کار و یا اینکه چه مقدار کار با موفقیت انجام شده است نیستند. از طرف دیگر افرادی معتقدند که چون ترجمه خودکار شکسپیر، گوته، تولستوی و دیگر ادیبان میسر نیست پس هیچگونه جایی برای هیچ نوع ترجمه مبتنی بر کامپیوتر وجود ندارد. آنها قادر نیستند سهمی را که ترجمه هایی نه کامل می توانند در کارشان یا در پیشرفت کلی ارتباطات بین المللی داشته باشند ارزیابی نمایند (5). در سالهای 1950 و 1960 رشته ترجمه ماشینی حوضه تحقیقی مهمی در زبانشناسی رایانه ای به حساب می آمد.

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در شروع، هدف مورد انتظار ترجمه خودکار انواع متون و اسناد در کیفیتی معادل یا حتی بهتر از مترجمان انسانی بود. اما خیلی زود معلوم شد که این هدف در آینده‌ای قابل پیش‌بینی غیر ممکن است. اگر نتایج ترجمه بخواهد به شکلی قابل چاپ در آید، اصلاح و تجدید نظر برون‌داد ترجمه ماشینی بوسیله انسان ضروری است. در عین حال این نکته نیز قابل درک بود که برون‌داد ترجمه ماشینی به شکل اصلاح نشده و خام نیز می‌تواند برای بسیاری از مقاصد از جمله برای خواندن متن به منظور پی بردن به ایده و محتوای کلی آن به یک زبان ناآشنا مفید واقع گردد. اما این کاربرد دوم ترجمه ماشینی برای سالها مورد غفلت واقع گشت (6).

بیشترین حجم ترجمه در دنیا از متونی نیست که از موقعیت ادبی و فرهنگی بالایی برخوردار باشند. اکثریت قریب به اتفاق مترجمان حرفه‌ای دست‌اندرکار ترجمه اسناد و متون علمی و فنی، معاملات تجاری و بازرگانی، اساسنامه‌های مدیریتی، اسناد حقوقی، دستورالعملها، کتابهای کشاورزی و پزشکی، پروانه‌های صنعتی، نشریه‌های تبلیغاتی، گزارشات روزنامه‌ای و غیره می‌باشند. مقداری از این کار پردرد سر و مشکل است. اما بیشتر آن کسل‌کننده و تکراری است ولی در عین حال نیاز به صحت و ثبات دارد. تقاضا برای چنین ترجمه‌هایی با سرعتی بیش از قابلیت حرفه مترجمی رو به افزایش است. مساعدت یک کامپیوتر جذابیت‌های واضح و ضروری دارد. سودمندی عملی یک سیستم ترجمه ماشینی نهایتاً بوسیله کیفیت برون‌داد آن تعیین می‌شود. اما آنچه که به عنوان یک ترجمه خوب صرف‌نظر از اینکه بوسیله کامپیوتر باشد یا انسان تلقی شود مفهومی است که تعریف دقیق آن مشکل به نظر می‌رسد. در عین حال بیشتر آن وابسته به موقعیتهای ویژه‌ای که در آنها ترجمه صورت می‌گیرد و نیز دریافت‌کننده‌های ویژه‌ای که ترجمه برای آنها انجام می‌شود است. صداقت، صحت، وضوح، سبک مناسب و سیاق همه از جمله ملاک‌هایی هستند که می‌توان از آنها استفاده کرد، اما باز هم اینها قضاوت‌های ذهنی هستند. تا آنجایی که به ترجمه ماشینی مربوط می‌شود آنچه که در عمل مهم است این است که چه مقدار تغییرات بایستی صورت گیرد تا برون‌داد ترجمه به حد استاندارد قابل قبولی از نظر مترجم یا خواننده انسانی رسانده شود. با چنین مفهوم دشوار و بی‌ثباتی چون ترجمه، محققان و سازندگان ترجمه ماشینی نهایتاً آرزو دارند که حداقل بتوانند ترجمه‌هایی تولید کنند که در موقعیتهای ویژه‌ای که مجبور به توضیح دادن اهداف تحقیق هستند مفید واقع شوند و یا کاربردهای مناسب دیگری برای ترجمه‌هایشان بیابند.

ترجمه ماشینی بخشی از حوزه وسیعتر "تحقیق محض" در پردازش زبان طبیعی مبتنی بر کامپیوتر در زبانشناسی

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

رایانه ای و هوش مصنوعی می باشد که مکانیزم های اساسی زبان و ذهن را بوسیله مدلسازی و شبیه سازی در برنامه های کامپیوتری مورد بررسی قرار می دهد. تحقیق در زمینه ترجمه ماشینی با پذیرش و بکارگیری جنبه های نظری و روشهای عملی در فرآیندهای ترجمه شدیداً با این امور سروکار دارد و به نوبه خود می تواند بینش ها و راه حلهایی را برای مشکلات و مسائل خاص خود بیابد. علاوه بر این ترجمه ماشینی می تواند بستر آزمایشی مناسبی را در مقیاس بزرگتر برای نظریه ها و روشهایی که بوسیله آزمایشات و تحقیقات در مقیاس کوچک در زبانشناسی رایانه ای و هوش مصنوعی انجام می شود فراهم نماید.

مانع اصلی بر سر راه ترجمه بوسیله کامپیوتر مانع کامپیوتری یا محاسبات نیست بلکه مانع زبانی یا زبانشناختی می باشد. مشکلات زبانی عبارتند از: ابهام واژگانی، پیچیدگی نحوی، تفاوت های واژگانی بین زبانها، ساختارهای غیر دستوری و مستتر و به طور خلاصه استخراج معنای جملات و متون از تجزیه نشانه های نوشتاری و تولید جملات و متون به مجموعه دیگری از علامات زبانی با معنای معادل. در نتیجه ترجمه ماشینی به شدت متکی به پیشرفتهای تحقیقات زبانی به ویژه شاخه هایی از آن با درجات بالایی از فرمالیته بودن است و در حقیقت همیشه همینطور بوده و خواهد بود. اما ترجمه ماشینی نمی تواند نظریه های زبانی را مستقیماً به کار گیرد. زبانشناسان با شرح مکانیزم های زیر ساختی تولید و درک زبان سروکار دارند. آنها بر روی مشخصه های قطعی متمرکزند و نه کوشش برای توصیف یا توضیح در مورد هر چیزی. در مقابل، ترجمه ماشینی با متون واقعی سروکار دارد. سیستم های ترجمه ماشینی با دامنه وسیعی از پدیده های زبانی، پیچیدگی های اصطلاحات، غلط های املائی، واژه های جدید و جنبه های اجرایی که همیشه هم در رابطه با زبانشناسی نظری مجرد نیستند سروکار دارد.

بطور خلاصه، ترجمه ماشینی به خودی خود یک شاخه مستقلی از تحقیق محض نیست بلکه هرگونه ایده، روش و فنی را که ممکن است در توسعه سیستم های پیشرفته سهمیم باشد خواه از زبانشناسی یا علم کامپیوتر، هوش مصنوعی و یا نظریه ترجمه باشد، اقتباس میکند. در واقع ترجمه ماشینی تحقیق کاربردی است، اما شاخه ای که از حجم مهمی از فنون و مفاهیم که به نوبه خود می توانند در حوضه های دیگر پردازش زبان مبتنی بر کامپیوتر به کار روند تشکیل شده است (5).

2-1- تاریخچه مختصر ترجمه ماشینی

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

استفاده از واژه نامه های خودکار برای غلبه بر موانع زبانی اولین بار در قرن هفدهم پیشنهاد شد. در قرنهای بعد پیشنهادهای (طرحهای) زیادی برای ایجاد زبانهای بین المللی (مانند اسپرانتو¹ به عنوان مشهورترین آنها) ارائه شد لیکن طرحهای اندکی برای مکانیزه کردن ترجمه تا اواسط قرن 20 ارائه گردید. در سال 1933 دو ابداع به طور جداگانه در فرانسه و روسیه به ثبت رسید. یک فرانسوی - ارمنی بنام جورج آرتسرونی² یک وسیله ذخیره ای را بر روی نوار کاغذی طراحی کرده بود که می توانست برای پیدا کردن معادل هر واژه به زبان دیگر مورد استفاده قرار گیرد. یک نسخه اصلی از این کار ظاهراً در سال 1937 عرضه شد. طرحی که به وسیله پتر اسمیرنف ترویانسکی³ روسی ارائه شد در آن زمان مهم تر به نظر می رسید و سه مرحله برای ترجمه مکانیکی در نظر می گرفت. در مرحله اول ویرایشگری که فقط زبان مبدأ را می دانست تجزیه منطقی واژه ها به صورت اشکال اصلی آنها و نقشهای نحوی را انجام میداد. در مرحله دوم ماشین می بایستی زنجیره هایی از اشکال اصلی و نقشها را به زنجیره های معادل آنها در زبان مقصد تبدیل کند. و در نهایت ویرایشگر دیگری که فقط زبان مقصد را می دانست این برونداد را به اشکال عادی آن زبان تبدیل می نمود. ترویانسکی از زمان خودش جلو بوده و خارج از روسیه ناشناخته باقی مانده بود، حتی زمانی که امکان استفاده از کامپیوترها برای ترجمه اولین بار به وسیله وارن ویور⁴ و آندره دی بوت⁵ مورد بحث قرار گرفت. اما این ویور بود که در جولای 1949 برای اولین بار توجه عموم را به ایده ایده ترجمه ماشینی جلب کرده و روشهایی از قبیل: استفاده از فنون رمز نویسی⁶ دوران جنگ، تجزیه آماری، نظریه اطلاعاتی شنون⁷، و جستجوی (کشف) مشخصه های زیربنایی منطقی و جهانی زبان را پیشنهاد نمود. ولی در مجموع لئون دوسترت⁸ در دانشگاه جورج تاون با همکاری IBM اولین پروژه عملی را که نتیجه آن اولین نمایش عمومی سیستم ترجمه ماشینی در ژانویه 1954 بود به ثمر رسانید. در این پروژه نمونه های به دقت انتخاب شده از جمله های روسی با استفاده از واژگان بسیار محدود در حدود 250 واژه و فقط شش قاعده دستوری به انگلیسی ترجمه شد. اگر چه این پروژه ارزش علمی کمی داشت و چندان موفق از آب در نیامد اما برای سرمایه

¹ - Esperanto

² -Gorge Artsrouni

³ - Petr Smirnov-Troyanskii

⁴ - Warren Weaver

⁵ - Andrew D.Booth

⁶ - cryptography

⁷ - shannon's information theory

⁸ -Leon Dostert

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

گذاری وسیع در زمینه تحقیقات ترجمه ماشینی در ایالات متحده محرک بسیار مؤثری بوده و مشوقی برای شروع پروژه‌های ترجمه ماشینی در هر کجای دیگر دنیا به ویژه در شوروی سابق بود (5).

بطور خلاصه می‌توان گفت که در این دوره ماشین‌های ترجمه نسل اول تولید شدند. این نسل از ماشین‌های ترجمه اغلب وابسته به فرهنگ لغت¹ بوده و تحلیل‌های گرامری سطح پایین داشتند و استفاده از خصیصه‌های معنایی از جمله ویژگی‌های آنها بوده است.

در دهه بعد ترجمه ماشینی رشد بیشتری نمود. در این دوره ماشین‌های ترجمه نسل دوم توسعه یافتند که این نسل از ماشین‌ها مبتنی بر قواعد نوشتاری² بوده و نمایش خلاصه‌ای از متن تبدیل یافته ارائه می‌نمودند و اغلب از رویکرد تبدیل استفاده می‌کردند. در این دوره که به دوره کمال گرایی³ مشهور است تأکید تحقیقات بر روی جستجوی نظریه‌ها و روش‌هایی برای رسیدن به ترجمه‌های کامل بود و نه چیزی کمتر از آن، که این خود تا حد زیادی مخرب بود چون هیچ سیستم کاملاً خودکار که قادر به ترجمه‌هایی با کیفیت بالا باشد وجود نداشته است و شاید هنوز هم وجود نداشته باشد (7).

در یک ارزیابی از ترجمه ماشینی که در سال 1960 انجام شد بار-هیلر⁴ بر این فرضیه متداول که هدف تحقیقات ترجمه ماشینی باید تولید ترجمه کاملاً خودکار با کیفیت بالا بوده و نتایج ترجمه‌های آنها بایستی غیر قابل تشخیص از نتایج ترجمه‌های انسانی باشد خرده گرفت. او چنین بحث کرد که موانع معنایی بر سر راه ترجمه ماشینی اساساً می‌توانند تنها به وسیله وارد کردن مقدار زیادی دانش دایره المعارفی در مورد جهان واقعی برداشته شوند. توصیه او این بود که ترجمه ماشینی نباید زیاد بلند پروازی کند، و بایستی سیستم‌هایی را بسازد که استفاده مقرون به صرفه از تعامل بین انسان و ماشین نماید. (نظریه‌ای که او برای اولین بار تقریباً ده سال قبل از آن یعنی زمانی که ترجمه ماشینی هنوز در طفولیت بود ارائه کرد). در سال 1966 گزارشی توسط شرکت‌های دولتی ترجمه ماشینی در ایالات متحده بنام ALPAC منتشر شد مبنی بر اینکه ترجمه ماشینی کندتر، ناصحیح‌تر و دو برابر گرانتر از ترجمه انسانی است و اینکه هیچ‌امیدی به ساخت ترجمه ماشینی مفید نمی‌رود. آنها دیگر نیازی نمی‌

¹ -Dictionary Driven

² -Rule Based

³ - perfectionism

⁴ -Bar-Hiller

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

دیدند تا بر روی تحقیقات ترجمه ماشینی سرمایه گذاری کنند و به جای آن توصیه کردند که کمک های ماشینی ای مانند فرهنگهای ماشین خوان برای مترجمان و حمایت مداوم از تحقیقات اساسی در زمینه زبانشناسی رایانه ای انجام شود. تأثیر گزارش ALPAC بسیار زیاد و وسیع بود هر چند که این انتقاد درست نبود. انتشار و نفوذ این گزارش آنقدر عمیق بود که خط پایانی برای تحقیقات ترجمه ماشینی برای بیش از یک دهه در ایالات متحده بود و برای سالها بعد از آن احساس عموم را نسبت به ترجمه ماشینی خدشه دار نمود (5).

در دهه بعد تحقیقات زیادی در زمینه ترجمه ماشینی خارج از ایالات متحده انجام گرفت و منجر به ساخت سیستم هایی مانند سیستم میتو¹ در مونترال برای ترجمه گزارشات آب و هوایی در رسانه های عمومی در سال 1976 و ساخت سیستم یوروترا²، یک سیستم چند زبانه براساس آخرین پیشرفتهای ترجمه ماشینی و زبانشناسی رایانه ای در سال های پایانی 1970 گشت (6).

سیستم ترجمه ماشینی در اواخر سالهای 1970 بالاخره به بازار تجارت راه یافت. یعنی زمانی که شرکت زیراکس کورپ³ از این سیستم برای ترجمه هزاران صفحه از اسناد رایج در طی یک سال استفاده کرد در واقع به محصولات چند زبانی اجازه ورود به بازارها و نمایشگاههای جهانی را داد. در ابتدا کاربران در اندیشه اینکه این نرم افزار می تواند ترجمه های کامل را تولید نماید آن را خریداری می کردند لیکن امروزه همانطور که می دانیم ترجمه ماشینی بیشتر به عنوان یک ابزار برای ترجمه کامل بکار گرفته می شود تا یک مترجم کامل (8).

تا قبل از سالهای 1980 سیستم های ترجمه ماشینی عمدتاً بر اساس ترجمه مستقیم بوسیله فرهنگهای دو زبانه بودند که از تجزیه نحوی نسبتاً پایینی نیز برخوردار بودند. اما با پیشرفت در زمینه زبانشناسی رایانه ای روشهای مناسب تری در این زمینه اتخاذ شد و برخی سیستم ها از روشهای غیر مستقیم نیز در ترجمه بهره جستند. بدین ترتیب که متنهای زبان مبدأ به باز نمود مجردی از معنا تبدیل می شد و در پی آن برنامه هایی برای تشخیص ساختار واژه (صرف) و ساخت جمله (نحو) و نیز برای حل مسئله ابهام ساخته می شد (5).

در خلال سالهای 1980 طرح مبتنی بر انتقال با روشهای جدید نظریه چند زبانی در هم آمیخت. مهمترین آنها تحقیق بر روی سیستم های مبتنی بر دانش به خصوص در دانشگاه ملون و پیتزبورگ بود که براساس ساخت سیستم

¹ - Meteo system

² - Eurotra system

³ - Xerox Corp

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

های درک زبان طبیعی در محدوده جامعه هوش مصنوعی بوده است. بحث بر سر این است که ترجمه ماشینی بایستی فراتر از اطلاعات زبانی محض (نحو و معنا شناسی) رود؛ ترجمه با درک مفهوم متن سروکار دارد و بایستی به دانش جهان واقعی رجوع نماید.

مهمترین توسعه قرن بیستم ظهور سیستم های ترجمه ماشینی تجاری بود. سیستم هایی از قبیل لوگوس¹ آمریکایی و متال² ژاپنی از جمله مهمترین آنها بودند. تقریباً تمام اینگونه سیستم های اجرایی به شدت وابسته به عمل پس ویرایش برای تولید ترجمه قابل قبول بودند. اما پیش ویرایش نیز متداول بود. بدین ترتیب که در بعضی از سیستم ها حضور یک اپراتور هنگام ورود متن برای مشخص کردن مرزهای واژه ها یا حتی نشان دادن محدوده عبارات و بندها لازم بود.

تا اواخر سالهای 1980 سیستم های با قابلیت ترجمه حجم زیادی از متون و کیفیت پائین عرضه شدند که اغلب آنها نیازمند پس ویرایشهای اساسی بودند. به منظور بهبود کیفیت ترجمه ماشینی تمهیدات زیادی در زمینه واژگان و نحو و نیز سبک متن قبل از ورود صورت گرفت و بدین ترتیب سیستم هایی چون سیستم³، لوگوس، متال و سیستم های مشابه بوسیله سازمانها و شرکتهای مهم بین المللی بکار گرفته شدند (10). احیای تحقیقات ترجمه ماشینی در سالهای 1980 و ظهور سیستم های ترجمه ماشینی در بازار تجارت سبب آگاهی عمومی روزافزونی در جهت اهمیت ابزار ترجمه شد. پیشرفتهای بیشتر در فناوری کامپیوتر، هوش مصنوعی و زبانشناسی نظری خطوط تحقیقات آینده را تعیین می نماید. اما اساسی ترین مشکل ترجمه مبتنی بر کامپیوتر مشکل آن با فناوری نیست بلکه با زبان، معنا، درک و تفاوتهای فرهنگی و اجتماعی جوامع بشری می باشد (5).

در واقع آغاز سالهای 1990 پیشرفتهای چشمگیری را در ترجمه ماشینی همراه با تغییرات اساسی در استراتژی از ترجمه مبتنی بر قوانین دستوری به ترجمه مبتنی بر پیکره های مبتنی بر نمونه ها (مثلاً برنامه ریورسو⁴) به خود دید و زبان، دیگر به عنوان یک وجود ساکنی که با قوانین ثابت و تغییر ناپذیر اداره می شد تلقی نمی گشت بلکه به عنوان یک پیکره فعال و متحرک که بسته به کاربرد و کاربر در طول زمان و بر طبق واقعیتهای فرهنگی و اجتماعی

1 - LOGOS

2 - METAL

3 - SYSTRAN

4- Reverso Program

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 رئوسرای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

تغییر پیدا می نمود در نظر گرفته می شد.

در این دوره نسل سوم ماشین های ترجمه بوجود آمدند. در این نسل ماشین های ترجمه از ترکیب رویکرد مبتنی بر قوانین نوشتاری نسل دوم با متد های مبتنی بر پیکره^۱ بوجود آمده بودند. ویژگی های این نسل شامل تحلیلی دستوری در سطح وابستگی ارتباطی تحلیل معنایی در سطح تشخیص جملات و اطلاعات لغوی از فرهنگ لغات بدست آمده و همراه با مشخصه های دستوری در زبان مبدا نگهداری می شوند ، می باشد. در این نسل متدهای مبتنی بر نمونه و روشهای آماری بکار گرفته شدند ، همچنین دانش جمع آوری شده محدود به حوزه مشخصی گردید و از بازخورد^۲ کاربران ماشین هم در جمع آوری آماری اطلاعات استفاده میشد . بطور خلاصه می توان گفت که اولین پیاده سازی ماشین های ترجمه مبتنی بر دانش^۳ در این دوره بوده است.

در دوره حاضر همچنان بهره برداری تجاری ماشین های ترجمه نسل دوم مشهود بوده و آگاهی نسبت به این تکنولوژی افزایش یافته است بطوریکه به متدهای ارزیابی ماشین های ترجمه نیاز مبرم می باشد که اگر این نیاز مرتفع نگردد به زودی کلیه تلاشها در جهت پیشرفت این تکنولوژی با خطر بی اعتمادی روبرو می گردند و طراحی ها در این دوره ، جهت تحقق اهدافی مانند ترکیب ماشین های ترجمه با سیستم های حافظه ترجمه ، ماشین های ترجمه سخنگو و سیستم های همه منظوره و بالا بردن کیفیت ترجمه ها می باشند

امروزه بسیاری از سیستم های تجاری ترجمه ماشینی از حافظه ترجمه بهره می جویند و بسیاری از سیستم های حافظه ترجمه نیز به نوبه خود با روشهای جدیدتر ترجمه ماشینی پیشرفته تر می شوند. در واقع مشخصه عمده سالهای 1990 افزایش سریع استفاده از ترجمه ماشینی و ابزار ترجمه است.

تا امروز ترجمه ماشینی همچنان به پیشرفت خود ادامه داده است. شرکتهای بزرگ بیشتر از آن استفاده کرده و فروش نرم افزارهای آن در بازار رو به افزایش است. این موقعیت سبب ایجاد سرویسهای ترجمه ماشینی شبکه ای (متصل به اینترنت) مانند آلتاویستا^۴ که سرویسهای پست الکترونیکی سریع، صفحات وب و غیره را به زبان دلخواه ارائه میدهند شده است.(49)

1- Corpus_Based

2 - Feed back

3 - Knowledge-Base

4 - Altavista

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

3-1- انواع سیستم های ترجمه ماشینی

فرض عموم بر این است که سیستم های ترجمه ماشینی به صورت یک جعبه سیاه جداگانه و مجزا وجود دارند. اما اکثر سیستم های ترجمه ماشینی خود از دو بخش عمده موتور ترجمه و یک یا چند واژه نامه تشکیل شده اند که طرق عملکرد این بخشها از سیستمی به سیستم دیگر متفاوت است. همچنین روشهای مختلفی که در ترجمه ماشینی اتخاذ می شود عمدتاً بر اساس منابع دانش مختلفی است که در کار ترجمه مورد استفاده قرار می گیرد.

1-3-1- سیستم های مستقیم

سیستم های قدیمی اغلب از روش سیستم مستقیم که به سیستم های ترجمه ماشینی نسل اول معروف بودند استفاده می کردند. این روش از ابتدائی ترین روشهاست که از یک فرآیند یک مرحله ای استفاده نموده و جملات زبان مبدأ با توجه کمی به تفاوت های ترتیب واژه ها به طور مستقیم و بدون مراحل واسطه ای به جملات زبان مقصد تبدیل می شوند. این روش بر اساس فرهنگ لغات غنی و تجزیه صرفی تنظیم شده است (9).

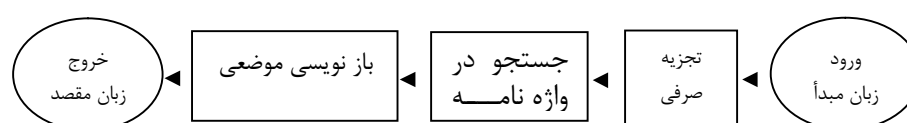
در رابطه با عملکرد سیستم های ترجمه ماشینی نسل اول باید گفت که کامپیوترهای موجود در اواخر سالهای 1950 و اوایل 1960 حتی در مقایسه با محقر ترین محاسبه گرهای الکترونیکی امروزی بسیار ابتدایی بودند. آنها فاقد زبانهای برنامه نویسی سطح بالا بوده و اکثراً برنامه نویسی ها بوسیله کد اسمبلی انجام می شد. به طور کلی سیستم های ترجمه ماشینی مستقیم نسل اول با مرحله ای شبیه آنچه که ما امروزه تجزیه صرفی می شناسیم آغاز می شد که در این مرحله تشخیص پایانه ها و تبدیل و تقلیل اشکال صرفی¹ مختلف به اشکال ستاکی غیر صرفی² انجام شده و نتیجه به عنوان ورودی برنامه جستجوگر در واژه نامه دو زبانه عمل می کرد. هیچگونه تجزیه ای از ساخت نحوی و یا روابط معنایی وجود نداشت. به عبارت دیگر تشخیص واژگانی تنها بر اساس تجزیه صرفی بوده و مستقیماً به واژه نامه دو زبانه منتهی می شد که در آنجا معادل های واژگانی زبان مقصد مشخص می گشت. البته برخی قواعد بازنویسی برای دستیابی به برون داد قابل قبول تر وجود داشت (مثلاً قواعدی که صفتها یا اجزای فعل را

¹ - inflected

² - uninflected

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

جابجا می کردند.) و پس از آن بلافاصله متن زبان مقصد تولید می شد. روش مستقیم را به طور خلاصه می توان به صورت شکل 1 نمایش داد.



شکل 1. سیستم ترجمه ماشینی مستقیم

محدودیت‌های بسیار زیاد این روش واضح است. این روش به عنوان ترجمه واژه به واژه با برخی انطباقات ترتیب واژه ای موضعی¹ شناخته می شود. کیفیت ترجمه آن همانند ترجمه شخصی است که دارای واژه نامه دو زبانه بسیار ابتدایی و ساده ترین نوع دانش دستور زبان مقصد بوده و اشتباهات ترجمه ای مکرری در سطح واژگانی داشته و ساختهای نحوی بسیار نامناسب در زبان مقصد که منعکس کننده ساختهای نحوی زبان مبدأ است را تولید می کند (10).

روش مستقیم روشی یکطرفه است یعنی اگر زبان مقصد بخواهد دوباره به زبان مبدأ ترجمه شود یک مبدل متفاوت دیگری باید مورد استفاده قرار گیرد. طبیعی است که برای این روش معایبی هم می توان بر شمرد. اولاً بدون یک تجزیه دستوری پیچیده تر ترتیب واژه ها در جمله زبان مقصد اغلب اشتباه خواهد بود. علاوه بر این در صورت وجود ابهام واژگانی در متن مبدأ، ترجمه نا درست حاصل خواهد شد. تجزیه رابطه بین قسمت‌های مختلف جمله اغلب وجود ندارد و بنابراین منجر به ترجمه ضعیف یا نامفهوم خواهد شد. روشهای جستجوی دو زبانه خام اغلب منجر به اشتباهاتی مانند ترجمه اشتباه معروف روسی "می تر بوام میرا" (ما جهان را تقاضا می کنیم) به جای "ما صلح می خواهیم" و بسیاری دیگر از اشتباهات می شود. روشهای مشابه که از قوانین ترتیب واژه ها نیز پیروی می کنند اگر چه ممکن است ترتیب اسم و صفت را رعایت کنند تا عبارت فرانسوی un chat noir (یک گربه

¹-Local

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

سیاه) را درست ترجمه نماید اما زمانی که با عبارت "یک گربه خیلی سیاه" روبرو می‌شود ترجمه نادرست un tres chat noir را به جای un chat tres noir ارائه می‌نماید. بنابراین روشهای مستقیم در مواجهه با ابهام واژگانی و حتی تفاوت‌های نسبتاً ساده در نحو با شکست مواجه می‌شوند. سادگی زبانشناختی و کامپیوتری این روش به سرعت شناخته شد. از نقطه نظر زبانشناختی آنچه که این روش فاقد آن بود تجزیه ساخت داخلی متن مبدأ به ویژه روابط دستوری بین قسمتهای اصلی جمله بود. فقدان پیچیدگی‌های کامپیوتری هم عمدتاً انعکاسی از ابتدایی بودن علم کامپیوتر در آن زمان بود (10).

2-3-1- سیستم‌های غیر مستقیم

شکست سیستم‌های ترجمه ماشینی نسل اول منجر به ظهور روشهای زبانشناختی پیشرفته تری برای ترجمه گشت. در واقع این روشها که به روشهای غیر مستقیم معروفند متن زبان مبدأ را با واسطه یک بازنمود میانجی به زبان مقصد تبدیل می‌کنند. این بازنمود میانجی نمایی از معناست که می‌تواند اساس تولید متن مقصد را تشکیل دهد (5). به گفته هارولد سومرز¹ این ممکن است به عنوان بازنمودی از معنای متن باشد یا حتی با ابهام کمتر، نمایی از ساخت نحوی متن باشد (11).

دو گونه اصلی روش غیرمستقیم عبارتند از: روش انتقالی و روش میان زبان، که در ذیل به توضیح آنها می‌پردازیم.

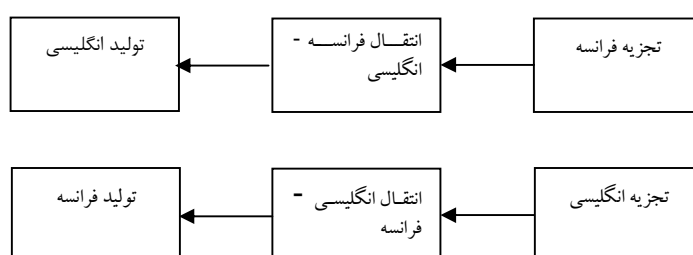
1-3-2-1 روش انتقالی

اولین گونه از روش غیرمستقیم روش انتقالی نامیده می‌شود. در واقع تمام سیستم‌های ترجمه درگیر نوعی انتقال هستند یعنی تبدیل متن یا نمایش مبدأ به متن یا نمایش مقصد. اصطلاح "روش انتقالی" به سیستم‌هایی اطلاق می‌شود که بخش‌های دو زبانی را بین بازنمودهای میانجی قرار می‌دهند اما این بازنمودها وابسته به زبان هستند. نتیجه تجزیه، بازنمود مجردی از متن مبدأ است و درونداد، تولید بازنمود مجردی از متن مقصد می‌باشد. عملکرد بخش‌های انتقالی دو زبانه تبدیل بازنمودهای (میانجی) زبان مبدأ به بازنمودهای (میانجی) زبان مقصد می‌باشد.

¹ - Harold Somers

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

شکل 2 این عملکرد را نشان می‌دهد. از آنجایی که این بازنمود ها بخش های مجزا (تجزیه، انتقال، تولید) را به هم ربط می دهند به آنها بازنمود های اتصالی نیز گفته می شود.



شکل 2. مدل انتقالی با دو جفت زبان

در روش انتقالی نحو زبان مبدأ - مشابه آنچه که قبلاً (به طور سنتی) بوسیله زبان‌شناسان انجام می‌گرفت - تجزیه می شود. سپس درخت نحوی به وجود آمده تبدیل به ساختی می‌شود که مناسب با زبان مقصد باشد. در نتیجه ترجمه ای در قالب ساخت زبان مقصد تولید می شود. یکی از معایب این روش نسبت به روش میان زبان این است که اضافه نمودن یک زبان جدید به سیستم نه تنها مستلزم دو بخش تجزیه و تولید است بلکه اضافه کردن بخش های انتقالی جدیدی را نیز می طلبد که تعداد آنها بسته به تعداد زبانهای موجود در سیستم تفاوت می کند. مثلاً در یک سیستم دو زبانه، زبان سوم نیاز به چهار بخش انتقالی جدید دارد. اما از طرف دیگر پیچیدگی نسبی بخش های تجزیه و تولید در یک سیستم انتقالی کمتر از روش میان زبان است چرا که بازنمود های میانجی در گیر در سیستم در روش انتقالی وابسته به زبان بوده و مشخصه هایی از متون مبدأ و مقصد می‌باشند.

اکثر سیستم های مبتنی بر انتقال، ترجمه را به کمک بازنمود میانجی انجام می دهند. اما نکته مهم در اینجا این است که این بازنمود ها وابسته به زبان هستند یعنی در آنها واژه های نمونه زبان مبدأ و مقصد مشاهده می شود. همچنین آنها ساخت زبانهای مبدأ و مقصد در گیر در ترجمه را، خواه به طور سطحی (صوری) خواه با عمق بیشتر، منعکس می‌نمایند، و این مشخصه یکی از تفاوت‌های عمده روش انتقالی با روش میان زبان (که در بخش بعدی به معرفی آن می پردازیم) است. بدین ترتیب که سیستم های میان زبان سعی در ارائه بازنمود های میانجی دارند که مستقل از زبان باشند چه از نظر واژگان و چه از نظر ساخت (5).

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در سیستم های مبتنی بر انتقال اساساً دو نوع مشکل وجود دارد، مشکلات واژگانی و مشکلات ساختی که در سطور زیر به بررسی آنها می پردازیم.

۱-۱-۲-۳-۱- انتقال واژگانی

انتقال واژگانی به معنای جایگزینی واحد واژگانی زبان مبدأ به وسیله واژه زبان مقصد است. واضح است که اگر تنها یک معادل برای واژه زبان مبدأ در زبان مقصد وجود داشته باشد مشکل خاصی به جز در مورد تفاوتی در ساخت نحوی وجود نخواهد داشت. لیکن این حالت تنها در حوضه ترجمه فنی که چنین تناظر واژگانی یک به یک وجود دارد مصداق پیدا می نماید. البته انتقال واژگانی در مورد ترجمه های تناظر چند به یک نیز مشکل زا نیست (مثلاً واژه های انگلیسی milk و lion که هر دو در فارسی به معنای شیر هستند)، بلکه مشکل از انتقال واژگانی تناظرهای یک به چند ناشی می شود. (مثلاً واژه انگلیسی term که در فارسی به معنای دوره، اصطلاح و واژه می باشد). در برخی موارد این امکان برای سیستم ترجمه ماشینی وجود دارد تا ترجمه های خاصی را اغماض کند و این به خاطر مهجور شدگی آن واژه ها یا به ندرت مورد استفاده قرار گرفتن آنهاست و یا می تواند به این دلیل باشد که واژه خاصی تنها در برخی از گونه های زبانی خاص دیده می شود. استفاده از واژه آکبند به جای واژه فارسی ناسوده در ترجمه ها مثالی از این مورد است. حتی واژه کامپیوتر به جای رایانه در بسیاری از موارد می تواند مورد استفاده قرار گیرد چرا که استفاده از واژه کامپیوتر برای فارسی زبانان راحت تر و جا افتاده تر از واژه رایانه است.

جایی که چند ترجمه مختلف برای یک واژه وجود داشته باشد طبیعی است که انتقال واژگانی نیاز به بررسی محیط یا بافت اطراف آن واژه پیدا می کند. مثلاً واژه انگلیسی know بسته به مفعول خود می تواند به شناختن یا دانستن یا معادل های دیگر در فارسی ترجمه شود. بدین ترتیب که ابتدا بایستی مفعول این فعل شناسایی شده و مشخصه های معنایی آن از قبیل جاندار یا بی جان بودن آن نیز تعریف گردد، سپس با استفاده از این بافت ترجمه

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع ملی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

مناسب این فعل انجام شود. مثال دیگر می‌تواند واژه انگلیسی morphology باشد که بسته به موضوع متنی که واژه در آن یافت می‌شود ترجمه می‌شود. در متون پزشکی این واژه به معنای ریخت شناسی است در حالیکه در متون مربوط به زبان‌شناسی و دستور زبان از آن با نام صرف یاد می‌کنیم. یافتن ترجمه مناسب برای چنین واژه‌هایی (که با توجه به موضوع متن تعیین می‌شوند) برای کامپیوتر بسیار مشکل تر از مورد نخست به نظر می‌رسد هر چند برای ترجمه انسانی تشخیص موضوع متن کار ساده‌ایست (12).

2-1-2-3-1- انتقال ساختاری

انتقال ساختاری زمانی ضروری به نظر می‌رسد که ساختاری که زبان مبدأ دارد برای زبان مقصد نامناسب باشد. برخی از مشکلات انتقال ساختاری نسبت به برخی دیگر راحت تر برطرف می‌شوند به طوریکه می‌توان برای انتقال آنها از زبان مبدأ به زبان مقصد قواعدی تعریف کرد و به صورت فرمول در ترجمه از آنها استفاده نمود. برای مثال به جمله انگلیسی زیر و ترجمه فارسی آن توجه کنید.

- Ali sent his friend a present.

- علی برای دوستش یک هدیه فرستاد.

با توجه به ساختار SOV در زبان فارسی و SVO در زبان انگلیسی تشخیص فاعل، فعل و مفعول در جمله زبان مبدأ و تبدیل آن به ساختار مناسب در زبان مقصد چندان مشکل به نظر نمی‌رسد. لیکن مشکل اصلی از آنجایی ناشی می‌شود که تفاوت‌های ساختاری زبانهای مبدأ و مقصد عمیق تر باشد. تبدیل جمله انگلیسی مجهول به جمله فارسی معلوم با همان مفهوم به دلیل قابل قبول نبودن بعضی ساختارهای مجهول در زبان فارسی می‌تواند مثال خوبی در این زمینه باشد (برای اطلاعات بیشتر رجوع شود به 13).

هم روش مستقیم و هم روش انتقالی در ترجمه ماشینی تکیه بر برنامه ریزی جداگانه ای برای هر جفت زبان درگیر که در جهت مشخصی ترجمه می‌شوند دارند. این مشکل با اضافه شدن جفت زبانهای بیشتری به سرعت

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

بیشتر می‌شود. روش انتقالی در n زبان درگیر در سیستم ترجمه از n^2 واحد انتقالی¹، n واحد تجزیه، و n واحد ترکیب استفاده می‌نماید. بنابراین یکی از معایب آن تعداد بسیار بالای قوانینی است که برای انجام آن استفاده می‌شود. مسئله دیگر حالت تحت اللفظی بودن ترجمه حاصل است. اگر چه روش انتقالی سبک واژه به واژه روش مستقیم را دنبال نمی‌کند اما پردازش معنایی کمی از متن مورد ترجمه را انجام می‌دهد. بنابراین با احتمال زیاد نمی‌تواند از عهده زبان عامیانه و مصطلح بر آید (14).

۲-۳-۱- روش زبان میانجی^۲

اخیراً محققان کوشش کرده‌اند تا معماری ای ارائه نمایند که در آن متون زبان مبدأ بوسیله نوعی بازنمود معنایی مستقل از زبان، زیربنایی و خنثی (بی طرف) به زبان مقصد تبدیل شود (برای اطلاعات بیشتر نگاه کنید به آرنولد و دیگران 1994). این کار از نظر تئوری تأثیر زبان مبدأ بر روی متن برون داد (در زبان مقصد) را از بین می‌برد. در یک سیستم میان زبان صرف هیچ نوع قانون انتقالی وجود ندارد چرا که بازنمود های میانجی بایستی در همه زبانهای درگیر در سیستم مشترک باشند.

جدا بودن فرآیندهای تجزیه و تولید در این روش یکی از معایب آن محسوب می‌شود. بدین مفهوم که این روش برای هر زبان مبدأ یک تجزیه‌گر و برای هر زبان مقصد یک تولیدکننده نیاز دارد. تجزیه متن مبدأ نیاز به یک تجزیه معنایی عمیقی که خود نیز مستلزم داشتن دانش گسترده جهان بیرون است دارد. متأسفانه معنای واقعی یک جمله همیشه به دست نمی‌آید و علاوه بر این اگر هم یک متن به همان عمقی که تصور می‌شود تجزیه شود باز هم قسمت بیشتر سبک نویسنده زبان مبدأ القا نمی‌شود (15).

روش میان زبان مخصوصاً برای سیستم های چند زبانه کاربردی تر است و زبانهای مقصد هیچگونه تأثیری بر روی هیچ یک از فرآیندهای تجزیه ندارند. مزیت این روش در این است که اضافه کردن یک زبان جدید به سیستم مستلزم ایجاد تنها دو بخش جدید است: دستور زبان تجزیه و دستور زبان تولید، که در این حالت ترجمه از زبانهای دیگر به زبان جدید و از زبان جدید به زبانهای دیگر مقدور می‌باشد (5). لازم به ذکر است که روش میان

¹ - Transfer modules

² - interlingua

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

زبان توجه کمی به فقدان دانش جهان بیرون^۱ که قبلاً به عنوان مانعی بر سر راه ترجمه ماشینی با کیفیت بالا تشخیص داده شده بود دارد.

اگر چه تنها سه روش اصلی در معماری ترجمه ماشینی وجود دارد که عبارتند از: روش مستقیم، روش انتقالی و روش میان زبان، شخص می تواند چندین بخش اطلاعاتی را اعمال نماید تا کارآیی سیستم ترجمه را افزایش دهد. چنین بخشهای اطلاعاتی شامل بخشهای مبتنی بر دانش، مبتنی بر آمار، مبتنی بر نمونه و غیره می باشند که در بخشهای بعدی به توصیف آنها خواهیم پرداخت.

3-3-1- روشهای مبتنی بر قانون^۲

در ماشین های ترجمه مبتنی بر قوانین ، پردازشی که صورت می گیرد ، تحلیل و نمایش معنی متن به زبان مبدا و تولید ترجمه معادل متن به زبان مقصد است . مسئله مورد بحث در اینجا روش بازنمایی است که باید از لحاظ ساختاری و لغوی کاملاً واضح و صریح باشد. این باز نمایی در روشهای مبتنی بر قانون ، یا روشهای انتقالی و میان زبان هستند و یا روشهای مبتنی بر دانش. روشهای انتقالی و میان زبان همچنانکه که گفته شد بر این اساس هستند که موفقیت در ترجمه ماشینی عملی مستلزم تعریفی از یک سطح بازنمودی برای متون است که در عین حال که مجرد است به حد کافی صوری هم هست تا بتواند در مورد هر دو زبان در گیر در ترجمه قابل انطباق باشد.

رویکرد تبدیل شامل سه مرحله می باشد ، ابتدا پردازش روی نمایش بدست آمده در زبان مبدا و سپس تبدیل آن به زبان مقصد و در آخر ترکیب و بازنمایی به زبان مقصد .

رویکرد میان زبان شامل دو مرحله میباشد ،مرحله اول تحلیل روی نمایش یک زبان واسط و سپس بازنمایی و تبدیل آن به زبان مقصد.

3-3-1- روشهای مبتنی بر دانش^۳

تاکنون بحث عمدتاً بر روی روشهای زبانشناختی مسائل انتقال در ترجمه بوده است، در حالیکه تفاوت‌های

¹ - world knowledge

² - rule-based approaches

³ - knowledge-based approaches

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

واژگانی ترجمه نیز به طور اجتناب ناپذیری انعکاسی از تفاوت‌های زمینه‌های فرهنگی و تفاوت‌هایی در اختلافات تصویری از واقعیت یکسان می‌باشند. بنابراین برخی معتقدند که ترجمه بایستی بر اساس بازنمود های تصویری غیر زبانی یا به بیان دیگر بر اساس معنایی که از فرآیندهای درک متن مشتق می‌شود باشد. درک یک متن یعنی برقرار کردن ارتباط بین آنچه که در متن گفته می‌شود، (محتوای زبانی آن) و پدیده‌های (فعالیتها، وقایع، ماهیت) خارج از متن (واقعیت غیر زبانی). خوانندگان متون را با توجه به آنچه که از جهان واقعی یا جهان‌هایی که خود می‌توانند تصور کنند می‌دانند تفسیر می‌کنند. اگر فرض شود که مفاهیم و درک برای تمام سخنگویان زبانهای بشری مشترک و جهانی هستند، نتیجه این می‌شود که این بازنمود های تصویری بین زبانی هستند و اینکه می‌توانند به صورت بازنمود های میانجی در سیستم های ترجمه ماشینی عمل می‌نمایند و نیز اینکه با کمک پایه های دانش مناسب، متون می‌توانند با این بازنمود های بین زبانی تجزیه شده و نیز از چنین بازنمود هایی تولید شوند. پس می‌توان گفت که استفاده از دانش جهان واقعی در سیستم های ترجمه ماشینی میان زبان کاربرد بیشتری دارد. یعنی به کمک دانش جهان واقعی نه تنها می‌توان ابهامات متن مبدأ را رفع نمود بلکه می‌توان بازنمود های مستقل از زبانی را به وجود آورد. مثلاً واژه انگلیسی wear دارای هشت واژه معادل در زبان ژاپنی با معانی ذیل می‌باشد:

Kitu (generic) , haoru (coat or jacket) , haku (shoes or trousers) , Kaburu (hat) , hameru (ring or gloves) , shimeru (belt or tie or scarf) , tsukeru (brooch or clip) , kakeru (glasses or necklace)

با توجه به این مثال پایه دانشی یک سیستم میان زبان باید قادر باشد تا شانزده (یا بیشتر) جنبه مختلف پوشیدن انواع لباس و اشیاء را از یکدیگر تشخیص دهد و این کار را مسلماً بر اساس دانشی که از انواع مختلف پوشیدنیها (مانند کلاه، کت، جوراب و غیره)، تفاوت‌های نژادی، فرهنگی، اجتماعی (لباس مراسم، لباس کار، لباس بچه، لباس گرمسیری، لباس سردسیری، لباس بارانی و غیره)، و فعالیت‌های مربوط به آنها (دوختن، خریدن، پوشیدن، درآوردن، تعمیر کردن و غیره) دارد باید انجام دهد.

اصطلاح ترجمه ماشینی مبتنی بر دانش به سیستم مبتنی بر قاعده‌ای (قانونی) اطلاق می‌شود که قادر به نشان

	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ

دادن دانش کاربردی و معنایی وسیعی در یک حوزه باشد که شامل توانایی تعقل در مورد مفاهیم آن حوزه است. مشخصه عمده روش مبتنی بر دانش تأکید زیاد آن بر روی درک کامل (از نظر کارکردی) معنای متن مبدأ قبل از ترجمه آن به متن مقصد است. اساساً قضیه به این شکل است که ترجمه با کیفیت بالا نیاز به درک عمیق متن دارد و ارائه این مدل حوزه ای^۱ برای این درک عمیق ضروری به نظر می‌رسد. یکی از نکات مهمی که در این رابطه بایستی به آن توجه نمود این است که پس ویرایش امری است وقت گیر و پرهزینه و بنابراین تلاشهایی برای تولید ترجمه با کیفیت بالا جریان طولانی ای را در بر خواهد گرفت. پس روش مبتنی بر دانش می‌تواند روشی باشد که درجاتی از درک متن را بر اساس حوزه تعریف شده ای از دانش در بر می‌گیرد. مدل حوزه ای که در این روش از آن استفاده می‌شود فقط تمام مفاهیم مربوط به یک حوزه خاص را نمایش می‌دهد و هیچگونه تعقل یا استنباط دیگری در مورد مفاهیم حوزه های دیگری که تعریف نشده اند بدست نمی‌دهد. نقش اساسی مدل حوزه ای رفع ابهام کامل از متن است و قسمت مهمی از آن مشخص نمودن محدودیتها یا قیدهایی که هر مفهوم در آن حوزه از نظر مفعولی که می‌گیرد و استدالات آن دارد است (مثلاً اینکه تنها موجودات زنده می‌میرند، تنها انسانها فکر می‌کنند و ...).

به محض وارد کردن دانش تعریف شده ای در ترجمه ماشینی با کیفیت بالا به وسیله درک متن شخص ترغیب می‌شود تا منابع دانش را مرتب اضافه نماید. کاملاً روشن است که حل پیش مرجع^۲ و حل دیگر ابهامات مرجعی نیاز به مراجعه به سطح ساختاری ای مافوق نحو و معنا شناسی جمله دارد. همینطور هم به دلایل سبکی برای افزایش پیوستگی متن، شخص نیاز به داشتن نمایش هایی از ساخت (ساختار) پاراگراف دارد. رسیدن به یک ترجمه با کیفیت واقعاً بالا به ویژه با انواع مختلف متون، ممکن است نیاز به شیوه مواجهه با استعاره^۳، کنایه^۴، کنش غیرمستقیم سخن^۵، نگرش سخنگو یا شنونده و غیره داشته باشد. در طول چند سال اخیر گروههای مختلفی در قسمتهای مختلف دنیا شروع به آزمایشاتی با الگو هایی که با دانش صریح یا اجزای قاعده ای که انواع بسیار متنوعی از اطلاعات را شامل می‌شدند نمودند. تمامی این روشها به هر شکلی که باشند ترجمه ماشینی مبتنی بر

¹ - domain model

² - anaphora resolution

³ - metaphor

⁴ - metonymy

⁵ - indirect speech act

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

دانش نامیده می شوند (17).

سیستم های مبتنی بر دانش ترجمه هایی با کیفیت بالا عرضه می‌دارند. با وجود این به خاطر نیاز به حجم زیاد دانشی که بتوان جملات زبانهای مختلف را به طور صحیح نمایش داد تولید این سیستم ها بسیار گران حاصل می شود. KANT و KANTOO سیستم های مبتنی بر دانشی در دانشگاه Carnegie meunon هستند که برای تولید اسناد چند زبانه (فرانسوی، اسپانیایی و انگلیسی) به کار می روند. این سیستم های میان زبانی شامل ابزار نگهداری واژگان و دانش هستند و ترجمه ها را تا حد امکان از نظر صحت بیمه نمایند.

پیچیدگی عملی این روش آن را به ناچار محدود به حوضه های بسیار محدودی با کاربردهای نسبتاً کم می نماید. همچنین لازم به ذکر است که سیستم های ترجمه ماشینی مبتنی بر دانش تاکنون در زمینه ابهام زدایی متون مبدأ متمرکز بوده اند و درگیر ایجاد بازنمود های مستقل از زبان جهانی برای فراهم کردن اطلاعات جهت تولید متن به زبان دیگری نیستند (5).

5-3-1- روشهای مبتنی بر پیکره (زبانی)¹

همچنان که دیدیم استفاده از روش مبتنی بر قاعده و مبتنی بر دانش در ترجمه ماشینی مشکلات عدیده ای را می آفریند. هیچ زبان طبیعی ای تماماً با مجموعه ای از قواعد مطابقت نمی کند و ابهام های زیادی به وجود می آید. علاوه بر این ترجمه ضرورتاً با بیش از یک زبان درگیر است که در اینصورت پیچیدگی افزایش می یابد. به این دلایل زبان شناسان شروع به جستجوی روشهای دیگری برای سیستم های ترجمه ماشینی نموده اند. برخی از این روشها از پیکره یا بدنه متن استفاده می کنند که به عنوان دریاچه ای از نمونه های آماده ای از کاربرد زبان عمل می کند. این روشها که به روشهای تجربی² معروفند اغلب از فنون آماری سطح پائینی استفاده می نمایند که یا به طور مستقیم به خود متون اعمال می شود و یا در تجزیه سطحی آنها به کار می رود. دلیل اینکه به این روشها

¹ - corpus-based approaches

² - empirical method

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

روشهای تجربی گفته می‌شود این است که در چنین روشهایی دانش زبانی ای که سیستم مورد استفاده قرار می‌دهد (از هر نوع آن که باشد) به طور تجربی از آزمایش متون واقعی استخراج می‌شود و نه از دانشی که به وسیله تعقل زبانشناسان به دست می‌آید.

این روشهای جدید هنوز مرحله کودکی شان را سپری می‌کنند. و این تا حدی به خاطر سخت افزار مورد نیاز برای ذخیره و دستیابی به حجم بسیار زیادی از داده های مورد نیاز در کار است که اخیراً شناسایی، فراهم و قابل دسترسی شده است. دلیل دیگر این است که ترجیح دادن دانش زبانی بر توانش زبانی که بوسیله چامسکی انجام شد و به عنوان اساسی برای تجزیه زبانی قرار گرفت مبین این نکته است که چنین کار تجربی ای به بهای توجه به روشهای نظری و صوری تر مورد غفلت واقع گشته است (14).

ماشین های ترجمه مبتنی بر پیکره به معنی هر نوع از ماشین های ترجمه مستقیم یا Direct می باشد که بر مبنای پیکره های زبانی موازی دو زبانه به صورت اتوماتیک عمل می کنند. منظور از ماشین های ترجمه مستقیم به هر تکنیکی که ترجمه ، در آنها مستقیماً از کلمات زبان مبدا به کلمات زبان مقصد، بدون استفاده از بازنمایی در یک ساختار زبانی واسط صورت می گیرد ، می باشد (16) .

در این نوع از ماشین های ترجمه از منبعی از اطلاعات شامل لغات ، ساختار گرامری و آمار و احتمالات که پیکره نام دارد جهت تحلیل، تبدیل و ترجمه استفاده می گردد .
در سطور زیر به معرفی برخی از این روشها می پردازیم:

1-3-5-1- روشهای مبتنی بر نمونه¹

پیشرفتهایی در زمینه فناوری کامپیوتر که دسترسی سریعتر به حافظه ها و ذخایر اطلاعاتی بیشتر را فراهم کرده است جستجوی روشهایی مبتنی بر دسترسی به پیکره های عظیمی از متون به زبانهای مبدا و مقصد را در رأس کار محققان قرار داده است. در این روش ترجمه اغلب با یافتن یا بازخوانی نمونه های مشابه و نیز با پیدا کردن و به خاطر سپاری اینکه چگونه عبارتی خاص از زبان مبدا یا چیزی شبیه به آن قبلاً ترجمه شده است سروکار دارد. این

¹ - example-based approaches

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

روش که به "ترجمه به کمک قیاس"¹ نیز معروف است روشی است که در آن ترجمه به وسیله مقایسه برونداد با پیکره ای از نمونه های ترجمه شده و استخراج نزدیکترین همانندها و استفاده از آنها به عنوان مدلی برای متن هدف (متن مقصد) انجام می شود (11).

پیشنهاد روشهای مبتنی بر نمونه به طور کلی به عنوان شقی دیگر در مقابل روشهای مبتنی بر دانش مطرح گردید و نیز مکملی برای روشهای تجزیه ای، انتقالی و تولیدی مبتنی بر قاعده به حساب می آید. بانکهای داده ای زبانی نمونه ها می توانند از تجزیه ساختاری پیکره عظیمی از متون مبدأ و ترجمه آنها به زبان مقصد که بوسیله مترجمان انسانی انجام شده است استخراج شوند. نتیجه مجموعه های دو زبانه ای از عبارات به زبان مقصد و ترجمه آنها به زبان مبدأ در بانک داده ای است (5). بانک داده ای متن از پیش تجزیه شده در حافظه ترجمه² ذخیره می شود. روشهای زیادی برای ذخیره کردن جفت‌های ترجمه ای در حافظه ترجمه وجود دارد. متن می تواند به صورت یک درخت تجزیه نحوی³ کامل یا به صورت جملات کامل و یا حتی عبارات ذخیره شود.

روش مبتنی بر نمونه به ویژه برای ترجمه عبارات اسمی پیچیده کاربرد بیشتری دارد. واژه های بسیاری در انگلیسی وجود دارد که دارای چندین معادل فارسی است. یافتن عباراتی که چنین واژه‌هایی در آنها قرار دارد و معادل فارسی این عبارات در پیکره زبان فارسی می تواند کمک خوبی در جهت ترجمه با ثبات چنین عبارات پیچیده ای باشد. حتی وجود مثالهای مشابه به عبارات اسمی در پیکره و محاسبه بعد (اندازه) شباهت احتمالی آنها می تواند ترجمه چنین عبارات مشکل زایی را آسان تر نماید. ترجمه چنین عباراتی با استفاده از نمونه های موجود در پیکره زبانی به مراتب راحت تر از متوسل شدن به قواعد مبتنی بر مقوله ها و ساختارهای دستوری و غیره می باشد. روشهای مبتنی بر نمونه همچنین می تواند در ترجمه جمله هایی که از نظر ساختاری مشابه جملات از قبل ترجمه شده هستند بکار رود. در این حالت نیز بعد شباهت نیاز است اگر چه این بار محاسبه این بعد بر اساس توزیع عناصر کلیدی مانند واژه های دستوری یا بر اساس زنجیره های خاصی از مقوله های دستوری یا ترکیبی از اینهاست (5). در توصیف خصوصیات این نوع ترجمه سومرز سه ملاک خاص را برای آن در نظر می گیرد.

الف) ترجمه ماشینی مبتنی بر نمونه از پیکره دو زبانه استفاده می کند.

1 - translation by analogy

2 - translation memory

3 - Parse tree

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

ب) ترجمه ماشینی مبتنی بر نمونه از پیکره دو زبانه به عنوان پایه دانش اصلی خود استفاده می نماید.

ج) ترجمه ماشینی مبتنی بر نمونه از پیکره دو زبانه به عنوان پایه دانش اصلی خود در زمان اجرای برنامه استفاده می نماید (18).

سومرز اشاره می نماید که دو ملاک نخست بسیار کلی هستند ولی ملاک سوم به اندازه کافی تعریف را محدود به ترجمه ماشینی مبتنی بر نمونه می نماید به طوریکه مثلاً آن را از ترجمه ماشینی مبتنی بر آمار که دو ملاک نخست نیز در آن صدق می کند اما احتمالات مستخرج از پیکره از قبل محاسبه می شوند متمایز می سازد (19).

این نکته نیز باید خاطر نشان شود که درجه عملی بودن این روش قطعاً بستگی به مجموعه داده های خوب دارد. با این وجود یکی از مزایای این روش این است که همچنان که مجموعه نمونه ها کامل تر می شود کیفیت ترجمه نیز بدون نیاز به بروز کردن و اصلاح توصیفات واژگانی و دستوری به طور فزاینده ای بهبود می نماید. بعلاوه این روش می تواند بسیار مؤثر واقع شود چرا که در بهترین حالت هیچگونه نیازی به اجرای کاربرد قوانین پیچیده نیست و تنها کاری که بایستی انجام شود پیدا کردن نمونه مناسب و (گاهی اوقات) محاسبه فاصله است. اما گاهی اوقات پیچیدگی هایی هم وجود دارد. مثلاً یک نوع مشکل زمانی به وجود می آید که یک واحد تعدادی نمونه متفاوت دارد که هر کدام از اینها با یک قسمت از زنجیره مطابقت می نماید ولی قسمتهایی که با آن واحد مطابقت می کنند همپوشی داشته و کل زنجیره را نمی پوشانند. در چنین حالاتی محاسبه بهترین تطابق نیاز به بررسی تعداد زیادی احتمالات دارد (17). و نکته آخر اینکه روش مبتنی بر نمونه می تواند در هر یک از مدل‌های اصلی: روش مستقیم، انتقالی یا میان زبان گنجانده شود. برخلاف روش مبتنی بر دانش که وابسته به تجزیه معنایی با درجه بالایی از تجرد است، این روش با تجزیه های روساختی نسبتاً سطحی و روشهای انتقال واژگانی ساده کار می کند، یعنی این روش می تواند به هر سطحی از انتقال، از تجزیه صرفی گرفته تا میان زبانی، اعمال شود (10).

2-5-3-1- ترجمه ماشینی آماری

طی چند سال اخیر توجه زیادی به استفاده از روشهای آماری در پردازش زبان طبیعی شده است. تجزیه داده

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

های آماری برای جمع آوری خودکار دانش لازم برای ترجمه ماشینی از متون دو زبانه متناظر¹ مورد استفاده قرار گرفته است. از نظر ترجمه ماشینی اصطلاح روشهای آماری به کاربرد فنون مبتنی بر آمار یا احتمالات در بخشهایی از کار ترجمه ماشینی (مانند ابهام زدایی معنایی واژه) اطلاق می شود. در سطور زیر سعی می شود تا روش مبتنی بر آمار محض در ترجمه ماشینی توصیف شود. این روش که برای اولین بار به وسیله براون معرفی شد (17) در بازشناسی گفتار² بسیار موفق بوده است که اگر چه نیاز به پیچیدگی های آماری زیادی دارد، اما تصور کلی آن کاملاً ساده به نظر می رسد. دو مفهوم کلیدی در این رابطه عبارتند از مدل زبانی³ و مدل ترجمه⁴. مدل زبانی احتمالات زنجیره واژه ای (در حقیقت جملات) را به ما می دهد که این احتمالات به وسیله $Pr(S)$ (برای یک جمله زبان مبدأ) و $Pr(T)$ (برای هر جمله زبان مبدأ مفروض) نشان داده می شوند. درواقع $Pr(S)$ احتمال وقوع زنجیره ای از واژه های زبان مبدأ و $Pr(T)$ احتمال وقوع زنجیره ای از واژه های زبان مقصد می باشد. مدل ترجمه نیز احتمالاتی را به ما می دهد. $Pr(T|S)$ که احتمال شرطی وقوع یک جمله زبان مقصد (T) در یک متن مقصد که ترجمه متنی است که شامل جمله زبان مبدأ (S) باشد است. حاصل این و احتمال وقوع S به تنهایی، یعنی $Pr(S) \times Pr(T|S)$ احتمال وقوع جفت جملات مبدأ - مقصد که به صورت $Pr(S, T)$ نوشته می شود را به دست می دهد. اما چگونه می شود احتمال وقوع یک زنجیره مبدأ (یا جمله مبدأ) یعنی $Pr(S)$ را محاسبه نمود؟ این احتمال را می شود به صورت احتمال واژه نخست ضرب در احتمالات شرطی واژه های متعاقب آن تجزیه نمود، یعنی:

$$Pr(S_3 | S_1, S_2), \text{ etc } \dots \times Pr(S_2 | S_1) \times Pr(S_1)$$

در واقع احتمال شرطی $Pr(S_2 | S_1)$ احتمال وقوع S_2 با دانستن اینکه S_1 رخ داده است می باشد. مثلاً احتمال اینکه am و are و در یک متن رخ دهند تقریباً یکسان است. اما احتمال اینکه am بعد از I رخ دهد بسیار بالاست در حالیکه احتمال اینکه are بعد از I رخ دهد بسیار پائین است. برای اینکه بتوان همه چیز را به صورت قابل کنترل و محدود در اختیار داشت معمول چنین است که تنها یک یا دو واژه قبل را در محاسبه این احتمالات

¹ - parallel bilingual texts

² - speech recognition

³ - language model

⁴ - translation model

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

شرطی در نظر می‌گیرند (که به ترتیب به مدل‌های گرام دو گانه و گرام سه گانه مشهورند). به منظور محاسبه این احتمالات زبان مبدأ (که مدل زبان مبدأ به وسیله تخمین پارامترها را تولید می‌کنند) به حجم بسیار زیادی داده یک زبانه نیاز می‌باشد چرا که اعتبار، کار آیی یا صحت مدل عمدتاً به حجم این پیکره داده ای بستگی خواهد داشت. جای دیگری که به حجم وسیعی از داده نیاز می‌باشد مشخص کردن پارامترهای مدل ترجمه می‌باشد که این نیاز به یک پیکره متناظر دو زبانه خواهد داشت. البته منابع نسبتاً کمی از چنین پیکره ای وجود دارد اگر چه گروه تحقیقی در IBM که عمدتاً مسئول توسعه چنین روشی بوده است موفق به تهیه سه میلیون جفت جملاتی از مذاکرات رسمی کانادایی¹ (فرانسوی - انگلیسی) که در آن هر جمله مبدأ متناظر با ترجمه اش در زبان مقصد می‌باشد به صورت یک پیکره متناظر از جملات شد (17).

پیکره های متناظر بر دو نوعند: پیکره های متناظر واژه ای و پیکره های متناظر جمله ای.

1-3-5-2-1- پیکره متناظر واژه ای

به مثال زیر توجه کنید :

1 2 3 4
a) It is perfectly clear.

1 2 3 4 5 6 7
b) I cannot find any mistake in it.

a`) Ka:melan(3) a:sheka:r(4) ast(2).

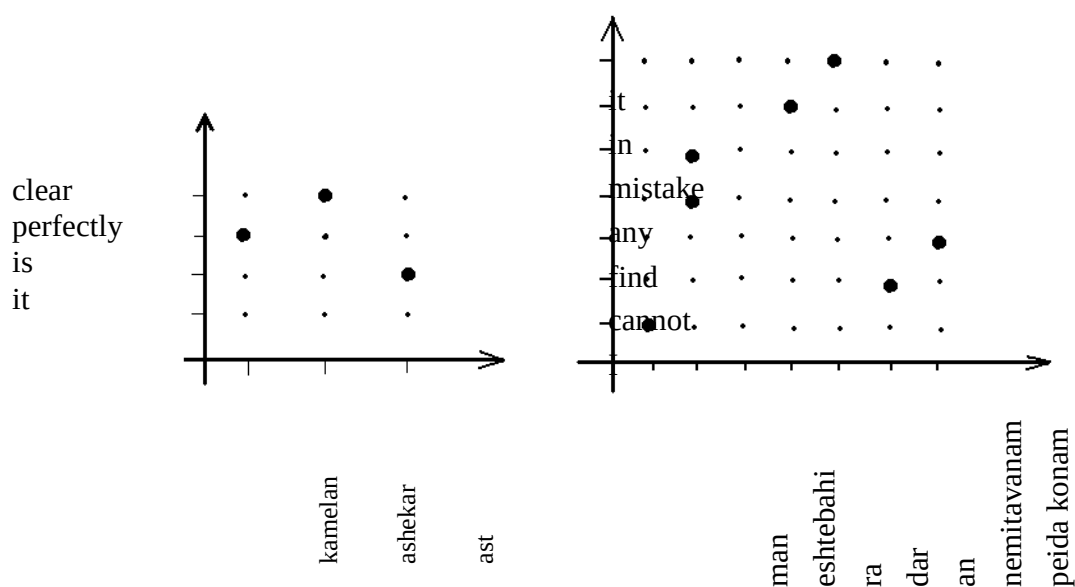
b`) man(1) eshteba:hi(4,5) ra: dar(6) a:n(7) nemitava:nam(2) Peida: konam(3)

در این مثال زبان مقصد، انگلیسی در نظر گرفته شده است. شماره های جلوی واژه های زبان مبدأ متناظر با ترتیب قرار گرفتن واژه ها در جمله متناظر زبان مقصد می‌باشند. اگر متناظر یک واژه در زبان مقصد وجود نداشته باشد طبیعتاً هیچ شماره ای هم در جلوی آن واژه زبان مبدأ قرار نمی‌گیرد (مثل واژه ra: در جمله فارسی)، و چنانچه بیش از یک واژه برای واژه متناظر مقصد یا مبدأ وجود داشته باشد آن نیز مشخص می‌شود (مثل واژه

¹ - Canadian Hansard

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع ماعن اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

eshteba:hi در جمله فارسی (سبک‌ها، متناظر، لایه شکا، دنگ، هم (شکا، 3) مر، ته ان نشان داد.



شکل 3. تناظر واژه به واژه

این نمودارها در بعضی از جفت زبانها تقریباً یکنواخت دیده می شود اما در زبانهای مانند انگلیسی و فارسی که عدم تناظر در برخی واژه ها و هم در ترتیب واژه ها مشاهده می شود یکنواخت نیست. این روشها بسیار امید بخش هستند از آن جهت که می توانند راه میان بری از جمع آوری دانشی باشد که همه کاربردهای پردازش زبان طبیعی را در بر می گیرد. اما متأسفانه روشهای ترجمه ماشینی آماری در جامعه ترجمه ماشینی به طور گسترده مورد استفاده قرار نگرفته است و این تا حدی به این دلیل است که ریاضیات درگیر در این روش برای محققان زبانشناسی کامپیوتری کاملاً و به طور دقیق جا نیفتاده است. دلیل دیگر هم این است که نرم افزارها و مجموعه داده های لازم در این روش به طور عموم در دسترس نیست. برای ایجاد زیر ساخت های نرم افزاری لازم جهت آزمایش و تحقیق در این زمینه کارهای زیادی بایستی انجام شود (12).

2-2-5-3-1- پیکره متناظر جمله ای

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی 2 — پ — 1	

نوع دوم پیکره های متناظر بر اساس جمله هستند که در آنها جمله متناظر هر جمله زبان مقصد به زبان مبدأ داده شده است.

روش مبتنی بر آمار اگر چه مزیت‌هایی دارد از جمله اینکه در آن نیاز به دانش زبانی به حداقل می رسد اما از نظر کاربرد دارای محدودیت‌هایی است چرا که بکارگیری این روش مستلزم فراهم کردن پیکره های حجیم یک زبانه و دو زبانه است که همانطور که می دانیم اکثر زبانها فاقد چنین پیکره هایی می باشند (12).

3-2-5-3-1- مراحل ساخت یک پیکره در ماشین ترجمه آماری

Ø جمع آوری متن ها یا لغات از هر دو زبان که این متون در واقع ترجمه معادل یکدیگر باشند و همچنین این متون باید از ارزش ادبی و کیفیت خوبی در ترجمه برخوردار باشند .

Ø پیش پردازش¹ و علامت گذاری متون

Ø هم تراز² و تعیین پارامترهای مدل و ساخت و آموزش مدل ها

Ø ارزیابی پیکره

الف - پیش پردازش و علامت گذاری متون

پیکره مجموعه ای متن به دو جفت زبان متفاوت می باشند این متن ها جمع آوری شده و معمولاً در مورد حوزه خاصی می باشند ولی ممکن است فرمت یکسانی نداشته و قابل خواندن توسط ماشین نباشند و همچنین دسته بندی نشده باشند . پیش پردازش شامل دسته بندی و یکسان سازی فرمت می باشد و سپس همه متن ها در کنار هم قرار می گیرند به طور مثال متن هایی که از آرشیو روزنامه بدست می آیند به فرمت HTML می باشند و شامل طبقه های³ مختلف هستند ابتدا این متون به فرمت متن¹ تبدیل می گردند و در دسته های مختلف تقسیم می شوند و

¹ Pre Processing

² Alignment

³ Category

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک-متن فارسی 2 — پ	

سپس همه در یک فایل بزرگ ذخیره می شوند . سپس روی متن های زبان مبدا و مقصد تفسیر² باید انجام گیرد و متن ها باید بخش بندی³ شوند یعنی بخش های مختلف مانند پاراگراف ، جمله و کلمه باید مشخص گردند و پس از بخش بندی ، برچسب گذاری⁴ باید روی جملات زبان مقصد و مبدا صورت گیرد و یا به بیان دیگر متن ها باید علامت گذاری⁵ شوند .

معمول ترین زبان نشانه گذاری SGML⁶ می باشد که امکان تعریف قاعده⁷ را برای متون تامین می نماید (48).

چند مثال متن علامت گذاری شده به صورت زیر قابل مشاهده می باشند :

<p><s>He left.</s>

<s></s></p>

<utt speak="Ali" Date="1385-07-01"> باید او باشد.</utt>

ب - هم ترازى

هم ترازى به معنی مشخص نمودن ارتباط بین واحد های دو زبان (واحد می تواند کلمات ، عبارات یا زیر کلمه باشد) ، در یک پیکره زبانی موازی است و همچنین مشخص نمودن اینکه چه کلماتی از هر زبان در یک موقعیت از ورودی ، به هم مرتبط هستند ، می باشد .

در این زمینه رویکرد های مختلفی ارائه شده است ولی دو رویکرد که قابل توجه می باشند به شرح زیر می باشند :

¹ Text

² Annotation

³ Segment

⁴ Tag

⁵ Mark

⁶ Standard Generalized Markup Language

⁷ Gramer

	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		شورای عالی اطلاع‌رسانی
	تاریخ: 1388/04/7	ویرایش: 1/0	
	کد زیرپروژه: پیک‌متن فارس 2 — پ		

رویکرد اول بر پایه استفاده از نرم افزار GIZA++ و پیاده سازی مدل مخفی مارکو و مدل IBM می باشد. با اینکه این رویکرد تنها به بخش هم تراز می پردازد و هم تراز به نوعی متغیر مخفی در مدل ترجمه¹ محسوب می گردد ولیکن نتایج بسیار خوبی حتی در پیکره های کوچک داشته است اما از دو نقص ساختاری برخوردار است که تاثیر شایانی در کارائی² این نوع سیستم ها دارند. یک مسئله اینکه یک عبارت چندین متناظر در زبان مقصد می تواند داشته باشد یا بالعکس ، یعنی هم تراز نامتقارن می باشد. این مشکل با روشهای متقارن سازی مختلفی که بر اساس ابتکارهای³ مختلف انجام میشوند ، که اکثرا از نظر زبانی نا آگاهانه می باشند در تحقیق برای یافتن استراتژی که سیستمهای posterior phrase-based با دقت بالا را تولید کنند برطرف شده است . و مسئله دیگر اینکه به دلیل پیچیدگی مدل های IBM و هزینه تخمین پارامترهای این مدل ، معرفی اطلاعات زبانی بسیار دشوار می باشد .

رویکرد دیگر بر اساس رخداد کلمات و احتمال اتصال کلمات به یکدیگر می باشد. سادگی و انعطاف پذیری در دریافت دانش و تقارن باعث شده که با وجود وابستگی به اطلاعات تجربی و استراتژی های tuning ، روشهای قابل توجهی محسوب گردند. بزرگترین عیب این روشها در اجبار آنها به هم تراز یک به یک است که منجر به هم تراز با دقت بالا و تعداد فراخوانی های کم می گردد و این عوامل باعث محدودیت در استفاده این روشها در ماشین های ترجمه عملی می شود.

تحقیقات سالهای اخیر بر روی ترجمه آماری بر بالا بردن کیفیت ترجمه با آموزش مدل ها که نه تنها بر مبنای لغات بلکه بر مبنای عبارات ، ترجمه صورت گیرد ، متمرکز بوده است .

در مدل های GIZA ابتدا هم تراز صورت می گیرد و سپس استراتژی های متقارن سازی اعمال می گردند تا به این ترتیب هسته اصلی هم تراز که عبارات از آن ساخته شده اند بدست آید . در اینجا با کمبود دانش زبانی و آموزش مدل ترجمه در یادگیری احتمالات از یک منبع نویزی⁴ روبرو هستیم .

در این رویکرد که معرفی میشود ، از روی تعداد تکرار رخداد کلمات در پیکره ، مستقیما به هم تراز عبارات⁵

¹Translation model

² Performance

³Heuristic

⁴Noize

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

عبارات¹ می‌رسیم، که بر این مشکل غلبه شده است. در اینجا قبل از اینکه تناظر بین کلمات را پیدا کنیم، پیوند² های با دقت بالا بین عبارات را پیدا می‌کنیم و جستجو برای اتصال بین عبارت ها، به مجموعه ای از عبارات ممکن محدود می‌گردد (48).

1-ب- اندازه گیری رخداد کلمات و عبارات

در این روش پارامتری به نام Dice-score برای هر جفت عبارت محاسبه میشود که میزان شایستگی هم ترازی بین دو عبارت یا کلمه را مشخص میکند که این مقدار از فرمول زیر بدست می‌آید:

$$\phi^2 = \frac{(ad - bc)^2}{(a + b)(a + c)(b + d)(c + d)}$$

رابطه 1-1- میزان شایستگی هم ترازی بین دو عبارت یا کلمه

a تعداد دفعاتی است که هر دو لغت یا عبارت با هم رخ داده اند، b و c تعداد دفعاتی است که یکی رخ داده و دیگری اتفاق نیفتاده و d تعداد دفعاتی که نه اولی و نه دومی در مجموعه داده ها³ رخ نداده اند، می‌باشند. در پیاده سازی مقدار باهم آیی⁴ یک کلمه یا عبارت که X بار در جمله به زبان مبدا و Y بار در جمله به زبان ترجمه اتفاق افتاده است، را برابر $\min(x, y)$ در نظر گرفته می‌شود.

با اینکه محاسبه این پارامتر ساده می‌باشد ولی در مواجهه با گرام دو گانه⁵ یا گرام سه گانه⁶ یا در واقع عبارات، بار محاسباتی ایجاد می‌نمایند ولی با این وجود در اکثر موارد اطلاعات قابل توجه و مفیدی را دارا می‌باشند. مشکل اصلی عملی نبودن محاسبه تمام ترکیبات، حتی اگر پیکره از تعداد لغات کمی تشکیل شده باشد، است.

¹ phrase-alignment

² Link

³ Data Set

⁴ Co-occurrence

⁵ Bi-gram

⁶ Tri-gram

	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	تاریخ: 1388/04/7	ویرایش: 1/0	
	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ		

چاره ای که برای رفع این مشکل در نظر گرفته شده از الگوریتمی استفاده می نمایم که تا حد امکان از اندازه گیری این باهم آیی ها اطلاعات مفید جمع آوری نموده و بدین ترتیب از بین تمام عبارات ممکن مجموعه ای را انتخاب نماید و از آن استفاده کند . برای این انتخاب از اطلاعات زبانی و آمار الگوریتم ، می تواند استفاده نماید (48).

2-ب - یک استراتژی در هم ترازی عبارات

این استراتژی شامل چهار مرحله می باشد :

ابتدا از مجموعه تمام کلمات ممکن ، مجموعه کوچکی از عبارت های جالب¹ را در هر زبان انتخاب می کنیم این انتخاب با کمک اطلاعات زبانی انجام میشود و مجموعه ای از عبارات را تشکیل می دهد که از چند کلمه تشکیل می شوند ولی نقش معنایی واحد را ایفا می کنند.

سپس الگوریتم با دقت بالا ، هم ترازی این عبارات را انجام می دهد و این عبارات را به عبارت دیگر یا کلمه پیوند می کند و از مقدار با هم آیی استفاده می نماید.

بعد از انجام هم ترازی در عبارات، الگوریتمی جهت هم ترازی واژه ها اعمال می گردد که مزیت وجود این مرحله کاهش پیچیدگی منتج شده از مرحله قبل می باشد.

در مرحله آخر پس پردازش روی هم ترازی بدست آمده ، اعمال می گردد که در این مرحله رفع ابهام روی مواردی مثل پیوندهای آزاد و همچنین تصمیم گیری در سطح جمله صورت می گیرد (48).

3-ب - طبقه بندی و انتخاب عبارات

¹ Interesting

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

همان طور که اشاره شد تمام ترکیبات ، رخداد هم زمان عبارات مختلف برای هر زبان بسیار زیاد می باشد و تقریباً غیر ممکن می گردد پس اولین کار این است که این فضای بزرگ را به بخش کوچکی که همان ترکیبات جالب از دیدهم تراز و ترجمه می باشند ، محدود می کنیم . "جالب" به عبارتی که به یک مفهوم اشاره دارند و یا از لحاظ معنایی به هم مرتبط می باشند گفته می شود و انتظار می رود که به یک کلمه در زبان مقصد هم تراز گردند یا اشاره نمایند و از طرفی طبقه بندی معنایی نمودن عبارات به اندازه گیری با هم آیی آنها کمک میکند(در یک مفهوم همان مقدار را به نمونه های مختلف نسبت می دهیم).

در انتخاب عبارات از اطلاعات زبانی کمک گرفته می شود و خصوصاً دو معیار برای انتخاب داریم که از دانش تکمیلی استفاده می کنند. ابتدا افعال را مشخص می نماییم که برای این کار از ماشین های خودکار قطعی که قوانین ساده را پیاده سازی می کنند استفاده می کنیم.

ورودی این قوانین، کلمات¹ و یا ریشه یا اصل²، می باشند و عبارت حاصل را به فعل اصلی،نگاشت³ می کنند به همین جهت طبقه بندی افعال در کاهش محاسبات در هم رخداد ها موثر است زیرا فرم عبارت هر چه باشد به فعل اصلی نگاشته می شود .

برای جلوگیری از اشتباه بعد از بدست آوردن ریشه یا فعل اصلی⁴ این فعل با دو لیست از افعال ممکن در زبان مورد نظر مقایسه می گردند (48).

یک نمونه از این قوانین در زبان انگلیسی به صورت زیر هستند :

¹ POS Tag

² lemma

³ Map

⁴ Head

 وزارت آموزش عالی و تحقیقات علمی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

PrP + VB VB(L=do) + PrP + VB VB(L=be) + PrP	PrP + MD(L=will/would/...) + VB MD(L=will/would/...) + PrP + VB
PrP + VB(L=be) + VBG VB(L=be) + PrP + VBG	PrP + VB(L=have) {+ been} + VB{G} VB(L=have) + PrP {+ been} + VB{G}

PrP: Personal Pronoun
VB / MD / VBG: Verb / Modal / Gerund (PennTree Bank POS)
L: Lemma (or base form)
{ } / (): optionality / instantiation

شکل 4. نمونه ای از قوانین طبقه بندی در زبان انگلیسی

معیار دیگر انتخاب اصطلاحات می باشد. مجموعه ای از عبارات مصطلح در دو زبان را داریم و قصد بر این است که یک کلمه معادل برای این عبارات پیدا نماییم و چون فرهنگ لغات مورد استفاده قرار نمی گیرد پس معیار انتخاب معنی نمی باشد.

عبارات دیگری که قابل طبقه بندی هستند اعداد و تاریخ یا زمان می باشد که مستقل از زبان بوده و در هر زبان قاعده خاص خود را دارند.

اگر اطلاعات زبانی در دسترس نباشند از اطلاعات آماری برای انتخاب مجموعه عبارات می توان استفاده کرد به طور مثال انتخاب بیشترین گرام دو گانه¹ یا گرام سه گانه یا به طور کل محتمل ترین ترکیب دو تایی یا سه تایی کلمات می تواند معیار انتخاب باشند.

باید توجه شود که در این مرحله تصمیم گرفته میشود که آیا اتصال روی یک عبارت صورت گیرد یا کلمات رها شوند تا در مرحله بعد به صورت کلمه به کلمه² متصل یا هم تراز گردند (48).

4-ب - هم تراز در سطح عبارات

در این مرحله میزان هم رخدادی بین یک عبارت از یک زبان در مقابل همه عبارات یا کلمات انتخابی از زبان دیگر محاسبه می گردد و یک استراتژی رقابتی برای متصل نمودن، استفاده می گردد که در این استراتژی یک

¹Most frequent bigram

² word-to-word

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

حد آستانه در نظر گرفته میشود و اتصال بین دو عبارت یا بین عبارت و یک کلمه بر اساس بهترین درجه با هم آیی و بالاتر بودن آن از حد آستانه انتخاب می گردد. این استراتژی بر مبنای این اصل است که عبارات در ترجمه مناسب تر از کلمات می باشند و حد آستانه مقداری است که محتمل ترین پیوند را تضمین می نماید و بدین ترتیب تمام اتصالات ممکن صورت نمی گیرد بلکه آنهایی که بالاترین ملاک هم رخدادی را داشته باشند انتخاب می شوند .

هنگامی که پیوند بین دو عبارت انتخاب می شود ، می توان با روشهایی بین کلمات دو عبارت نیز اتصال پیدا نمود ، اگر لزومی داشته باشد . این کار بوسیله الگوریتم های هم تراز کلمات می تواند انجام گردد در صورتی که محدود به جستجو در عبارات نیز باشد (48).

5-ب - هم تراز کلمات

از الگوریتم تکراری استفاده می شود که در آن ابتدا هم تراز اولیه با استفاده از با هم آیی کلمات انجام می گیرد و احتمال اتصالات محاسبه می شود سپس جستجو بر اساس اولین پیوند برتر¹ اعمال شده که ابتکار هایی که بکار می برد global aligned sentence link probability می باشد و جستجو در مراحل بعد با استفاده از اطلاعات گرامری نیز آگاهانه تر میشود و به دلیل بزرگی فضای جستجو استفاده از این ابتکار ها بسیار موثر است .

دو پارامتر که برای بهینه سازی کارایی الگوریتم استفاده می شوند

احتمال اینکه لغت به هیچ کلمه دیگری وصل نشود² که باعث قابل مقایسه بودن پاسخهای بطور کامل هم تراز شده³ و یا آنهایی که اصلاً هم تراز نشده اند⁴ و می گردد .
حداقل امتیازی⁵ که یک پیوند قابل قبول واقع شود .

همچنین محدود کردن جستجو ، در پیوندهای ممکن در زبان مقصد ، به یک پنجره که به این ترتیب کارایی الگوریتم افزایش می یابد و همچنین ابهام کاهش می یابد زیرا از وجود کلمات تکراری جلوگیری می کند .

¹Best first search

²Null probability

³Full aligned

⁴ Partialy aligned

⁵Minimum score

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

پنجره در همسایگی قطری که از تقسیم طول هر دو جمله بدست می آید، است. این پنجره وابسته به هر دو زبان بکار رفته بوده و ثابت نیست. برای انگلیسی اسپانیایی به تجربه مشخص شده که 8 کلمه بهینه است (48).

6-ب - پس پردازش

در این مرحله تصمیم برای هم ترازی نهایی گرفته میشود و از اطلاعات در سطح جمله استفاده می شود. کارهایی که انجام می گردد تصمیم راجع به کلمات پیوند نشده که آیا می تواند آنها را نیز پیوند کند یا خیر و جستجو برای پیوند های طولانی تر و بازبینی مجدد پیوند های کلمات و عبارات در تمام جفت جملات پیوند شده می باشند.

در حالت ایده آل این بخش به قسمتهای انتخاب عبارت و هم ترازی، بازخورد¹ می دهد که در تصمیمات قبلی که برای link ها در نظر گرفتند با اطلاعات کلی تجدید نظر نمایند (48).

6-3-1- ترجمه ماشینی مبتنی بر اصل²

سیستم های ترجمه ماشینی مبتنی بر اصل روشهای تجزیه ای را که بر اساس نظریه اصول و پارامترهای دستور زبان زایای چامسکی³ است به کار می گیرند. تجزیه گر نحوی یک ساخت نحوی کاملی که شامل اطلاعات واژگانی، عبارتی، دستوری و موضوعی باشد را تولید می نماید، و تأکید بیشتر بر روی توانمندی نمایش های مستقل از زبان و تجزیه های زبانی عمیق می باشد.

در ترجمه ماشینی مبتنی بر اصل دستور زبان به صورت یک مجموعه مستقل از زبان، اصول تعاملی خوش ساخت بودن و مجموعه ای از پارامترهای مستقل از زبان دیده می شود. بنابراین برای یک سیستمی که از Π زبان استفاده می کند، Π بخش پارامتری و یک بخش اصول لازم است. پس میتوان گفت این روش مناسب استفاده با معماری میان زبان می باشد. این روش برد وسیعی از پدیده های زبانی را در بر می گیرد، اما فاقد دانش عمیق

¹ Feed back

² - principle-based machine translation

³ - principle & parameters Theory of Chamsky's Generative Grammar .

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

درباره حوضه ترجمه که سیستم های مبتنی بر دانش و مبتنی بر نمونه بکار می برند می باشد (15).

3-3-1- روشهای ترکیبی

دهه اخیر شاهد ظهور تعدادی از روشهای جدید و امیدوارکننده ای برای ترجمه ماشینی بوده است و به تازگی تحقیقات در این زمینه متوجه ایجاد و توسعه روشهای ترکیبی برای طراحی سیستم های ترجمه ماشینی که بیش از یک روش را اتخاذ می کنند شده است. این روشهای ترکیبی مزیت هر دو یا چند روش را با یکدیگر دارند. پژوهشگران ترجمه ماشینی امروزه آزمایشات و تحقیقات متعددی را با ترکیبهای مختلفی از روشهای مختلف انجام داده تا به مزیتهای بیشتری دست یابند. تاکنون ترکیب سازی روشهای ترجمه ماشینی اساساً بر روی دو جنبه تأکید داشته است: جنبه فنی و کیفیت ترجمه پیشرفته. مثال قابل ملاحظه در این رابطه اجرای همزمان سه موتور ترجمه ماشینی مختلف و ترکیب سه برونداد آنهاست (7).

ایده کلی در روشهای ترکیبی استفاده از یک مدل زبانی برای تجزیه متن مبدأ و یک مدل غیر زبانی برای یافتن تفسیر مناسب است.

روش مبتنی بر نمونه در ترجمه جملاتی که قبلاً در حافظه ترجمه نمایش داده شده اند بسیار خوب عمل می نماید. روش مبتنی بر آمار قادر به تولید جملات با صحت بالا می باشد، اما عموماً در برخورد با اصطلاحات، با هم آیی ها و توابع منقطع (با فاصله طولانی دور از هم) خیلی موفق نیست. با ترکیب این دو روش یک سیستم ترکیبی مبتنی بر نمونه و مبتنی بر آمار بدست می آید که ابتدا در حافظه ترجمه جملات منطبق را جستجو می نماید، اگر هیچ انطباق نزدیکی پیدا نشد آنگاه تجزیه آماری و تفسیر جملات مورد استفاده قرار می گیرد.

به هر حال آزمایشات و نیز تجربیات در زمینه ترکیب سازی در روشهای ترجمه ماشینی به شدت رو به افزایش است. انواع ترکیب سازیهای انجام شده در این زمینه عبارتند از:

الف - ترکیب اطلاعات آماری در سیستمهای مبتنی بر قاعده با استفاده از تکنیکهای مختلف (21 و 20)

ب - استخراج واحدهای ترجمانی جدید از متون دوزبانه و تبدیل آنها به سیستم های ترجمه ماشینی مبتنی بر قاعده (22)،

ج - ترکیب حافظه های ترجمه با ترجمه ماشینی مبتنی بر قاعده (9)،

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

د - ترکیب ترجمه ماشینی مبتنی بر نمونه با ترجمه ماشینی مبتنی بر قاعده (23)،

ه - ترکیب حافظه ترجمه با ترجمه ماشینی مبتنی بر نمونه (24).

علیرغم وجود نتایج عملی و تنوع تکنیکهای استفاده شده، مایه بسی تعجب است که هنوز هیچ نظریه ترجمه ماشینی ترکیبی ارائه نشده است.

4-1- مکانیسم، مؤلفه‌ها و اجزای سازنده سیستم‌های ترجمه ماشینی

در ساختار سیستم‌های ترجمه ماشینی سه مؤلفه اصلی وجود دارد و هر مؤلفه نیز به نوبه خود از مجموعه اجزایی تشکیل می‌شود که در راستای تحقق نقش یا نقشهای آن سه مؤلفه اصلی عمل می‌کنند. مؤلفه‌های فوق‌الذکر در تناظر با سه سطح که علم زبان‌شناسی نوین برای هر زبان طبیعی و زنده قائل است قرار دارند. این سه سطح عبارتند از:

الف - واژگان

ب - نحو و ساختار

ج - معناشناسی

باید همواره این نکته را در نظر داشت که تفکیک این سطوح مبین عدم ارتباط بین مختصات واژگان، نحو و ساختار و معناشناسی متن نیست، بلکه هر سه سطح در سلسله مراتب چندبعدی خود در حکم تار و پود یکدیگرند. به عبارت ساده‌تر این تفکیک صوری و اعتباری است، و در راستای نوعی شبیه‌سازی فرایند ترجمه، در ذهن مترجم است، در واقع ترجمه ماشینی از طریق این شبیه‌سازی متحقق می‌شود. از نظر اهمیت، اولویت، سهولت پردازش و یا مسأله‌زایی آنها نیز نمی‌توان بین سه مؤلفه یاد شده تفاوتی قائل شد، بلکه هر سه آنها، از دیدگاههای یاد شده هم‌سطح‌اند. شاید در وهله اول تصور شود که در پردازش‌ها، به علت انتزاعی‌تر بودن مؤلفه معناشناسی، این سطح پیچیده‌تر و در نتیجه مسأله‌زاتر باشد، ولی باید توجه داشت که در عمل چنین نیست، هر سطح حتی سطح واژگان نیز مسائل و مشکلات خاص خود دارد چه ابهام‌ها در هر سطح می‌تواند ظاهر شود. اگر در هر سطح به گونه‌ای منطقی با مسائل و قضایا برخورد شود از میزان اغماض و پیچیدگی و ابهام‌زایی در سطوح بعدی کاسته می‌شود. دلیل امر وجود تعامل و ارتباط ظریف بین عناصر این سه سطح است. در سطور زیر سعی می‌شود تصویری کلی از مکانیسم و نقش هر یک از مؤلفه‌های مورد بحث ارائه شود.

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

1-4-1- واژگان

نخستین مرحله پردازش در ترجمه ماشینی در سطح واژگان انجام می‌گیرد. پردازش مورد نظر عبارتست از تجزیه و تحلیل مرفولوژیکی واژه‌های واحد ترجمه. در این مرحله عناصر موجود در سطح واحد ترجمه یک به یک تفکیک، پردازش و مقوله‌گذاری می‌شوند. واژگان هر زبان مختصات و دشواریهای خاص خود را دارد و بالقوه می‌توانند مسأله‌زا باشند. برخی از این مشکلات به زبانی خاص اختصاص ندارند؛ از جمله مصادیق این مضمون می‌توان به مثال‌های زیر توجه داشت: در زبان فارسی واژه‌های ساده و پرکاربرد “مرد” و “آرد”، در زبان انگلیسی واژه‌های “run” و “correct” و در زبان فرانسه واژه‌های “livre” و “cours” می‌توانند به مقوله‌های گرامری متعدد، تعلق داشته باشند. و این امر جدا و سوای از پدیده چند معنایی است که در هر زبان امری عادی تلقی می‌شود. عدم پردازش صحیح این گونه‌موارد، می‌تواند موجب مسائل و مشکلات لاینحل در مراحل و سطوح بعد می‌گردد. شایان ذکر آن‌که، زبان انگلیسی از این منظر جزو خطراترین زبانها محسوب می‌شود ولی باید همواره به خاطر داشت که هم‌نویسه‌ها تنها عامل خطا در امر پردازش واژگان نیستند، در بعضی زبانها مشکلات خاص و مختص آن زبان وجود دارد که در عرصه پردازش زبان مجال ظهور و بروز یافته، منشأ ابهام و خطا می‌شوند. انفصال و ناپیوستگی عناصر واحد ترجمه در زبان انگلیسی، امکان ایجاد خطا و انتساب آنها به سازه‌های غیر، از جمله عوامل ابهام‌زا و خطا در مرحله پردازش واژگان محسوب می‌شود بنابراین باید با لحاظ همه موارد مسأله‌آفرین، تدابیری اندیشید که در این سطح - به ظاهر ساده - میزان خطا به حداقل ممکن برسد و راه برای انجام پردازش و تجزیه و تحلیل‌های زبان‌شناختی بعدی، بهتر هموار گردد. در جهت نیل به این هدف در زبان انگلیسی اقدامات متعدد انجام گرفته است. از جمله در سال 1991 پروژه‌ای عظیم به مدیریت کنسرسیومی متشکل از چند دانشگاه و نهاد علمی و صنعتی، من جمله دانشگاه اکسفورد و دانشگاه لانکاستر، مرکز تحقیقات کتابخانه ملی بریتانیا، شورای علمی و مهندسی، آکادمی انگلیس، به منظور تدوین پیکره زبانی انگلیسی BNC و مقوله‌گذاری واژه‌های انگلیسی شروع و در سال 1996 خاتمه یافت. پیکره زبانی براساس صد میلیون واژه تدوین شد و از طریق آن امکان کددهی و مقوله‌گذاری فراهم آمد و در نهایت، حدود 98% مقوله‌گذاری صحیح میسر گشت. (25)

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:			 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیر پروژه:			
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن‌فارس — 2 — پ		

2-4-1- نحو و ساختار

تجزیه^۱ در اصطلاح زبان شناسی رایانه ای^۲ یعنی اجرای برنامه تجزیه و تحلیل نحوی متن و یا به عبارت کلی تر تجزیه و تحلیل گرامری جمله از طریق کامپیوتر. و تجزیه کننده^۳ به سیستمی اطلاق می شود که با استعانت از آن واحد ترجمه از منظر ساختار و نحو تجزیه و تحلیل می شود. به نحوی که کامپیوتر می تواند از طریق آن به سهولت عناصر تشکیل دهنده واحد ترجمه را بازشناسی کند، نقش آنها و روابط درونی و فیما بین آن عناصر و نوع قانون حاکم بر آنها را تشخیص دهد. در ترجمه ماشینی دستیابی به اطلاعات یاد شده ضروری و اجتناب ناپذیر است، چه علاوه بر برگردان عناصر معنایی و واژگان، ترجمه ساختار نحوی واحد ترجمه، از زبان مبدأ به مقصد نیز باید انجام شود. در واقع برگردان قالب و ساختار جمله از زبان مبدأ به زبان مقصد، شرط اصلی و اساسی ترجمه است، و این امر جز با تجزیه و تحلیل گرامری واحد ترجمه میسر نمی شود. انجام این مهم از طریق سیستم های تجزیه و تحلیلگر رایانه ای صورت می گیرد. گفته شده است که مهمترین رکن ماشین ترجمه تجزیه و تحلیلگر نحوی یعنی تجزیه کننده است. هر چند این طرز نگرش خالی از اغراق نیست ولی بیانگر اهمیت دومین مؤلفه اصلی در ساختار سیستم ترجمه ماشینی است. این اهمیت، انگیزه آن بوده که تاکنون صدها نوع تجزیه کننده در سطح جهان ابداع شود، تنوع مکانیسم های بکار رفته در ساختار تجزیه و تحلیلگرهای نحوی واقعا اعجاب انگیز است. تقریباً هیچ دانشگاهی و مؤسسه آموزش عالی در سطح جهان یافت نمی شود که دارای مرکز تحقیقات زبان شناسی یا رایانه ای باشد و حداقل یک تجزیه و تحلیلگر نحو در آن ابداع نشده باشد. نقش اصلی تجزیه و تحلیلگر نحوی در ساختار ترجمه ماشینی عبارت است از:

- 1- تشخیص دستوری بودن جمله
- 2- شناخت عناصر سازنده جمله و شیوه ترکیب آنها به منظور تشکیل واحدهای بزرگتر، عبارت، بند و ...
- 3- تعیین نقش هر یک از واحدها در جمله
- 4- شناخت دقیق روابط معنایی بین عناصر سازنده جمله

¹-Parsing

²- Computational linguistics

³-Parser

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

علاوه بر نقش اصلی، نقشهای فرعی - بنا به مقتضای کار و سلیقه ابداع کننده - نیز برای تجزیه و تحلیل گرنحوی در نظر گرفته می‌شود.

شیوه کار تجزیه کننده ها متفاوت است ولی معمولاً براساس دو نوع شیوه کارشان، به دودسته تقسیم‌بندی می‌شوند که عبارتند از: تجزیه کننده های نزولی و تجزیه کننده های صعودی با آنکه هیچکدام از این دو نوع تجزیه کننده ها کامل و بی نقص نیستند ولی معمولاً تجزیه کننده نزولی بیشتر مورد توجه بوده ، که خود نشانه موفقیت آمیز بودن، تجزیه و تحلیل نحوی است. (25)

3-4-1- معناسازی

مشکل چندمعنایی بودن واژه‌ها، شناخته شده‌ترین ویژگی زبانهای طبیعی است و این امر زائیده قانونی است که در زبان‌شناسی تحت عنوان اقتصاد کلامی¹ اشتهار یافته است.

مترجم به هنگام برخورد به واژه چند معنا، با تکیه بر توانش زبانی² خود و بافت و موضوع متن می‌تواند معادل واقعی واژه را تشخیص دهد و برگزیند. رایانه فاقد چنین توانشی است. بنابراین طراح ماشین ترجمه ناگزیر است مقدماتی فراهم سازد و در الگوریتم ، به تمهیداتی دست یازد که از طریق آن بتواند خلأ موجود را ترمیم کند. این تمهیدات باید از نخستین مراحل پردازش یعنی از مرحله واژگان و نحو و ساختار آغاز گردد و تا آخرین مرحله یعنی مرحله تجزیه و تحلیل معناساختی متن تداوم یابد. مثالهای زیر می‌تواند روشنگر مضمون باشد، کلمه "prime" در زبان انگلیسی یک مقوله دستوری، چهار مقوله معنایی و 23 معنا دارد. واژه "Suspect" نیز به سه مقوله دستوری (اسم و فعل و صفت) تعلق دارد و دارای شش مقوله معنایی و 21 معناست، "minister" نیز دو مقوله گرامری، شش مقوله معنایی و 11 معنادارد.

قاعده³ ماشین ترجمه می‌تواند 1280 معادل فارسی را برای تعبیر "suspect The prime" انتخاب کند. ولی اگر سیستم در پردازش واژگان و ساختارها دقیق عمل کرده باشد تعداد انتخابها به 32 و در نهایت اگر در مرحله پردازش معنا، منطقی با قضا یا برخورد کند و فرهنگ معناساختی تدوین شده براساس اصول زبان‌شناسی رایانه‌ای در اختیار داشته باشد، سیستم منحصراً تنها یک گزینه را انتخاب خواهد کرد. براساس همین محاسبه تعداد معادل‌های

¹ - Economie Langagiere

² - Competence

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

محتمل فارسی “prime minister” در سیستم‌های ماشین ترجمه عادی 176 خواهد بود. ولی اگر در مراحل قبلی، براساس معیارهای دقیق، پردازش‌ها صورت گرفته باشد، تعداد معادل‌ها به 112 و در نهایت به تنها انتخاب صحیح دست خواهد یافت. واژه سه حرفی “run” را در نظر می‌گیریم، این واژه دو مقوله دستوری و 58 مقوله معنایی و بیش از 140 معنا دارد. تنها، سیستمی قادر خواهد بود که معنای مناسب این واژه را برگزیند که با استعانت از الگوریتمی قوی و فرهنگ معنانشناختی جامع بتواند مؤلفه‌های معنانشناختی را در طول پردازش and و or کند. (25)

5-1- ابهام زدایی معنی واژه

ابهام واژگانی اشاره به حالتی دارد که یا یک واحد واژگانی متعلق به طبقات دستوری یا اجزاء کلام مختلف با معنایی مختلف باشد و یا یک واحد واژگانی بیش از یک معنی داشته و معنایی مختلف آن متعلق به یک طبقه دستوری باشند (26). در اینجا هدف از ابهام زدایی معنی واژه رفع ابهام از حالت دوم است یعنی موردی که در آن معنایی مختلف واژه متعلق به طبقه دستوری یکسانی باشند. ابهام زدایی معنی واژه اصطلاحی است که به استخراج معنی صحیح و مقتضی از واژه‌هایی که چند معنایی هستند اطلاق می‌شود و این کار غیرممکن نیست چرا که اگرچه واژه‌های ممکن است معنایی زیادی داشته باشند، طبیعتاً تنها یکی از آن معنایی مناسب بافتی که واژه مورد نظر در آن ظاهر می‌شود است. برای مثال واژه “bat” می‌تواند به معنای یک جانور شب‌زی، وسیله‌ای برای بازی بیسبال و یا چشمک زدن باشد. هدف ابهام زدایی معنی واژه تعیین یکی از این معناها بر اساس بافت زبانی که این واژه در آن ظاهر شده می‌باشد. در بسیاری از سیستم‌هایی که کار ابهام زدایی معنی واژه را به عهده دارند سیستم برچسب گذاری اجزاء کلام وجود دارد که مقوله دستوری واژه‌ها را مشخص می‌نماید. مثلاً در مورد واژه “bat” در صورتیکه سیستم برچسب گذاری اجزاء کلام تشخیص دهد که این واژه اسم است آنگاه سیستم ابهام زدایی معنی واژه کار خود را شروع کرده و از میان دو معنای اسمی این واژه یعنی جانور شب‌زی و وسیله برای بازی بیسبال یکی را انتخاب می‌نماید.

ابهام زدایی معنی واژه یک امر ضروری در ترجمه ماشینی بوده و کاربرد زیادی در آن دارد چرا که سیستم ترجمه ماشینی بایستی قادر باشد تا مثلاً معنی واژه “bank” را در یک متن انگلیسی تشخیص داده و تعیین نماید

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

که آیا آن را به “بانک” یا به “ساحل رودخانه” در فارسی ترجمه نماید. از طرف دیگر سیستم های بازیافت اطلاعات نیز بایستی تنها متون و اسنادی را که معنی مناسب واژه داده شده برای جستجو را در برمی گیرند جستجو کرده و نمایش دهند. مثلاً اگر واژه جستجو یا کلیدی “بانک” باشد پس سیستم می بایستی تمامی نمایش های مربوط به “ساحل رودخانه” را حذف نماید.

روشهای موجود و متداول برای ابهام زدایی معنی واژه به سه نوع تقسیم می شوند که عبارتند از: روشهای مبتنی بر قاعده، روشهای مبتنی بر پیکره یا مبتنی بر آمار و روشهای ترکیبی.

1-5-1- روشهای مبتنی بر دانش یا مبتنی بر قاعده

این روشها منابع دانش واژگانی را همراه با اطلاعاتی در مورد اولویتهای گزینشی، تداعی مفاهیم و طبقات دستوری بکار می گیرد و در قالب قواعدی در اختیار گرفته و بدین وسیله از بین معانی مختلف واژه یکی را برمی گزینند. چنین سیستم هایی نتایج خوبی را به دنبال دارند ولی ساختن قواعد گران تمام می شود و علاوه بر این قواعد موجود در سیستم ممکن نیست که بتوانند تمام موقعیتهای و بافتهای ممکن موجود را پوشش دهند (12).

1-5-2- استفاده از روشهای آماری در ابهام زدایی معنایی

مانند دیگر کاربردهای پردازش زبان طبیعی، در سیستم های ابهام زدایی معنی واژه مبتنی بر آمار نیز استفاده از پیکره های زبانی که به طور خودکار یا دستی حاشیه نویسی و برچسب گذاری شده اند رو به افزایش است. مدل های مبتنی بر آمار که از روشهای نظارتی برای نسبت دادن معانی به واژه های مبهم استفاده می نمایند نیاز به یک پیکره زبانی که از نظر معنای واژه ها برچسب گذاری شده است دارد. معنایی با بیشترین احتمال که براساس داده های پیکره ای محاسبه می شود به وسیله سیستم انتخاب می شود. برای کاربردهای چون ابهام زدایی معنی واژه حجم بسیار زیادی از متون لازم است تا این اطمینان حاصل شود که تمامی معانی مختلف یک واژه چند معنایی در پیکره زبانی وجود دارد. برای تعیین معنای صحیح یک واژه مبهم با استفاده از پیکره های زبانی دو روش ممکن وجود دارد: استفاده از اطلاعات توزیعی¹ و استفاده از واژه های بافتی¹. اطلاعات توزیعی توزیع فراوانی

¹ - Distributional information

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی — 2 — پ	

معناهای یک واژه است در حالیکه واژه‌های بافتی به محدوده واژه‌های طرف راست و طرف چپ یک واژه خاص اطلاق می‌شود.

خط مشی آماری، ما را از نیاز به الگوریتم‌های ماشینی و دانش زبان‌شناسی گسترده برای حل مشکلات زبان‌شناسی بی‌نیاز می‌کند.

از آنجایی که استفاده از روشهای آماری در ابهام‌زدایی معنی واژه بسیار مؤثر بوده و به سرعت در حال پیشرفت و فراگیر شدن است در بخش بعدی به معرفی چند سیستم مبتنی بر آمار که هدف آنها ابهام‌زدایی معنی واژه است اشاره می‌شود (12).

1-2-5-1- استفاده از گنجواژه^۲

در این روش می‌توان از مشخصات عمومی معنای یک کلمه در فرهنگ لغت استفاده نمود. در این روش، تا کنون از سه نوع اطلاعات استفاده شده است. نوع اول روش لسک است که مستقیماً از تعریف معنای کلمات در فرهنگ لغت استفاده می‌شود. نوع دوم که نشان می‌دهد چگونه از اطلاعات طبقه‌بندی شده موجود در فرهنگ لغت می‌توان طبقه‌بندی معنایی یک کلمه را با توجه به متنی که آن کلمه در آن وجود دارد، به دست آورد. (یارافسکی^۳ 1992) و سوم اطلاعاتی است که از ترجمه یک کلمه به کمک یک فرهنگ لغت دو زبانه به دست می‌آید و استفاده از معنای کلمه در زبان مقصد برای رفع ابهام از معنای کلمه. (27)

یاروفسکی^۴ سیستمی را پیشنهاد می‌کند که معنی واژه در هر نوع متن انگلیسی را با استفاده از مدل‌های آماری طبقات اصلی دستوری از گنجواژه Roget ابهام‌زدایی می‌نماید. معانی مختلف یک واژه به طور طبیعی با طبقات مختلف گنجواژه مطابقت دارند از این رو انتخاب محتمل‌ترین طبقه راه سودمندی است برای مجزاکردن و تشخیص معناهای مختلف. محتمل‌ترین طبقه دستوری یک واژه بوسیله جستجوی صد واژه دوروبر (مجاور) نمایندگان هر طبقه تعیین می‌شود. این واژه‌های نماینده تعیین شده و در طی محاسبه بوسیله امتحان با صد واژه مجاور هر نمونه از یک طبقه بارگذاری می‌شوند. این سیستم نیاز به پیکره زبانی مربی^۵ بزرگی (که نیازی به

² - Context words

³ - Thesaurus

³ - Yarowsky

⁴ - Yarowsky

⁵ - training corpus

	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	تاریخ: 1388/04/7	ویرایش: 1/0	
	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ		

برچسب گذاری آن نیست) و یک گنجواژه (یا هر نوع منبع داده ای دیگر که بتواند مجموعه ای از طبقات معنایی و نیز واژه ها در هر طبقه را تعریف نماید) دارد. هیچ نوع بخش سیستمی دیگری به جز تجزیه اجزاء کلام و استخراج شکل ریشه ای واژه ها مورد نیاز نیست. اطلاعاتی که از داده های مربی به دست می آید شامل فهرستی از واژه های نماینده برای هر طبقه دستوری و بار¹ آنهاست. تمام واژه های موجود به اشکال ریشه ای خود تقلیل می یابند (یعنی اطلاعات صرفی نادیده گرفته می شود) تا با افزایش تعداد رخدادها بتوان به آمار دقیق تری دست یافت.

روش محاسبه بار به صورت لگاریتم برجستگی واژه² برای آن مقوله تعریف می شود. برجستگی یعنی $Pr(w)$ ، به عبارت دیگر احتمال اینکه یک واژه در یک بافت از یک مقوله داده شده ظاهر شود تقسیم بر احتمال وقوع آن واژه در پیکره زبانی به طور کلی (برای واژه های مفید، برجستگی احتمال دارد که بزرگتر از یک بوده و بنابراین بار آن یا لگاریتم برجستگی مثبت خواهد بود). این سیستم از صد واژه مجاور واژه مورد نظر استفاده می کند. تمام واژه ها در بافت بدون توجه به مقوله دستوریشان مورد توجه قرار می گیرند. اطلاعات نحوی و صرفی دخالتی در کار نداشته و تمام واژه ها به شکل ریشه ای خود در می آیند. برای محاسبه محتمل ترین طبقه یک واژه مراحل زیر اجرا می شوند:

1-1-2-5-1- محدود سازی طبقه

تنها طبقاتی که شامل این واژه می شوند به عنوان احتمالات (طبقات احتمالی) در نظر گرفته می شوند. این کار به طور کلی مفید است اما در صورتیکه در گنجواژه از قلم افتادگی هایی وجود داشته باشد می تواند سبب بروز مشکلاتی نیز گردد. یاروفسکی مثالی از واژه "suit" می زند که معنای آن در بازی ورق (که به معنای خال است) در گنجواژه آورده نشده بود و بنابراین اشتباه شماره گذاری شده بود. وقتی که واژه به صورت طبقه محدود نشود (به طوریکه تمام طبقات موجود باشند) همه نمونه ها به طور صحیح نشانه گذاری خواهند شد.

1-1-2-5-2- جمع بارها

¹ - weight
² - word's salience

			
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	شورای عالی اطلاع رسانی		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

با در نظر گرفتن هر طبقه ممکن به صورت انفرادی صد واژه بافتی مجاور برای هر واژه نماینده آن طبقه جستجو می شوند. بار هر نماینده‌ای که پیدا شد به جمع کل اضافه می شود. (اگر برای یک نماینده بیش از یک بار پیدا شود بار آن نیز بیش از یک بار اضافه می شود). این نوعی تکمیل سازی قانون بیز¹ برای $Pr (Cat | w)$ می باشد که برای هر W مربوطه در بافت:

$$\underset{\Sigma}{\text{Argmax}} \quad \text{Log} \frac{Pr(w | Cat) Pr(Cat)}{Pr(w)}$$

W in context

رابطه 1-2- جمع بارها

احتمال واژه $Pr(w)$ می تواند حذف شود چون تغییری در نتیجه ماکزیمم سازی نخواهد داد. (البته اگر نتایج بخواهند به صورت بین واژه ای مقایسه شوند وجود آن ضروری است). احتمال پیشین برای تمام طبقات یعنی $Pr(Cat)$ مساوی فرض می شود از اینرو این نیز در سیستم مورد استفاده قرار نخواهد گرفت.

3-1-2-1-5-1- انتخاب محتمل ترین طبقه

طبقه ای که بالاترین مجموع بار را داشته باشد به عنوان طبقه واژه مورد نظر انتخاب می شود. این سیستم در مقایسه با تلاشهای قبلی خوب عمل می نماید مخصوصاً به این خاطر که این یک سیستم واقعی است، محدود به واژگان خاصی نمی شود و در قلمرو نسبتاً وسیعی کار می کند (یک دایره المعارف). برای ارزشیابی این سیستم از تعدادی واژه‌های مبهم که قبلاً برای سیستم های ابهام زدایی دیگر مورد استفاده قرار گرفته شده استفاده شد و صحتی بین 72% تا 99% بدست آمد (28).

3-1-5-3- استفاده از موصوف برای دستیابی به معنای صفت

¹ - Bayesian rule

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

جاستسن^۱ و کتز^۲ یک روش اصولاً زبانی را برای نوعی ابهام زدایی که در آن تنها از سرنخهای از نظر معنایی یا نحوی مرتبط استفاده می شود ارائه می دهند. در این روش از روشهای آماری برای سازماندهی و تجزیه داده ها در پیکره های مربی مورد استفاده قرار می گیرند اما مرحله ابهام زدایی از تکنیکهای آماری استفاده نمی نماید. سیستم واقعی ای که آنها توصیف می کنند صفتها را تنها با استفاده از اسم هایی که بوسیله آن صفتها تعریف می شوند ابهام زدایی می نماید. در ابتدا سیستم به طور کامل خودکار است یعنی با استفاده از تکنیکهای آماری آموزش داده می شوند تا موصوفهایی را که معنی صفت هایشان را به نوعی نشان می دهند جدا کنند. برای اینکه بتوان پوشش وسیعتری را در اختیار داشت، سیستم بعداً با استفاده از قواعدی که با دست کد گذاری شده اند و نه اطلاعات آماری تکمیل می شود. جاستسن و کتز تمام آزمایشات خود را بر روی پنج صفت انگلیسی مبهمی که بیشترین فراوانی را دارند یعنی صفتهای *hard, light, old, right* و *short* انجام دادند.

برای انجام آزمایش یک پیکره مربی ابهام زدایی شده لازم است. برای پنج صفت فوق این کار می تواند به طور خودکار با استفاده از خاصیت متضادهای معنایی خاص انجام پذیرد یعنی معانی مختلف واژه *old* می تواند بوسیله متضادهایش یعنی *young* و *new* تعریف شود. هنگامیکه یک صفت و یکی از متضادهایش در یک جمله ظاهر شوند معمولاً یک اسم (اسم یکسانی) را توصیف می کنند و صفت مبهم هم تقریباً همیشه معنی عکس متضادش را دارد. بنابراین پیکره مربی می تواند بوسیله انتخاب جملاتی از پیکره بزرگی که جفت های متضاد در آن ظاهر شده اند و تجزیه آنها به منظور تضمین اینکه صفتها به اسم یکسانی برمی گردند تولید شود. پس یک تجزیه گر نحوی نیز ضروریست بدین منظور و به منظور تعیین اینکه کدام اسم بوسیله صفتی که در جملاتی که در حال ابهام زدایی هستند تعریف می شود. تجزیه گر صرفی نیز لازم است تا اطلاعات صرفی را از اسمها حذف نماید. این سه جزء اصلی یعنی پیکره مربی، تجزیه گر نحوی و تجزیه گر صرفی بوسیله جاستسن و کتز به طور دستی شبیه سازی شدند.

نتیجه ای که از پیکره مربی بدست آمد مجموعه ای از اسمهای شاخص برای هر معنای جداگانه یک صفت بود بطوریکه وقتی این اسم ها بوسیله آن صفت تعریف می شوند آن صفت همان معنای خاص خود را نشان میداد.

¹ - Justeson
² - Katz

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

برای تعیین اینکه چه اسم‌هایی شاخص‌های یک معنی صفت باشند برای هر اسمی تعداد با هم آیی (با هم آیی هر معنای صفت با آن اسم) از زیر پیکره ابهام زدایی شده استخراج می‌شود. اطلاعات صرفی هم از اسم‌ها جدا می‌شود، از اینرو تعداد اسم‌های مفید موجود جهت استنتاج آماری افزایش می‌یابد. به جهت اینکه زیر پیکره ابهام زدایی شده انتخاب تصادفی‌ای از یک پیکره واقعی نبوده است (چرا که این زیر پیکره تنها جملاتی با جفتهای متضاد را شامل می‌شود) سوگیری آماری¹ در آن مشاهده می‌شود. جاستسن و کتر فرمولی برای تعیین سطح اطمینان می‌دهند که این سوگیری را مدنظر قرار می‌دهد و در آخر مقاله شان آن را به عنوان پیوست ارائه می‌دهند. اسم‌هایی به عنوان شاخص مورد استفاده قرار می‌گیرند که معنای ویژه‌ای با سطح اطمینان 95% از خود نشان می‌دهند.

برای ابهام زدایی یک صفت ابتدا جمله بایستی تجزیه شود و اسمی که آن صفت آن را تعریف می‌کند مشخص گردد. اطلاعات صرفی باید از این اسم جدا شوند. اگر صورت پایه اسمی که تعریف می‌شود (یا نوع اسم طبقاتی مانند اسامی افراد و اسامی اشیاء نیز استفاده می‌شوند) شاخصی برای یکی از معناهای آن صفت باشد در آن صورت آن معنا انتخاب می‌شود و در صورتیکه شاخصی برای هیچ یک از معناهای آن صفت نباشد هیچ معنایی به آن تعلق نمی‌گیرد.

به علت پوشش محدودی که استفاده از این روش دارد (حتی اگر به جای آن اسم‌هایی که از نظر آماری شاخص‌های مهمی به شمار می‌روند تمام اسم‌ها شاخص فرض شوند) جاستسن و کتر بعدها سیستم متفاوت دیگری را پیشنهاد دادند که از صفتها (نسبت‌ها) ی نحوی و معنایی به جای اسم‌های واقعی برای اسم‌های مورد نظر استفاده می‌نماید و نیز از قوانین دست‌نوشته شده‌ای برای تخصیص معنا براساس آن صفتها (نسبت‌ها) استفاده می‌نماید.

(مثلاً برای light نسبت not dark → colour + استفاده می‌شود).

این روش از داده‌های آماری استفاده نمی‌کند بنابراین به جزئیات آن در اینجا نمی‌پردازیم. این سیستم در مورد هر صفت بر اساس مجموعه‌ای از صد جمله به طور تصادفی انتخاب شده از کل پیکره زبانی ایکه شامل آن صفت می‌شود آزمایش شد. جملاتی که دارای صفتی که معنای آن بوسیله یک متضاد قابل توضیح نبود بودند

¹ - statistical bias

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

مستثنی شده و خارج شدند و این سیستم صحتی معادل 100% برای اسمهای شاخص مهم با سطح پوشش 22/6% و صحتی معادل 97% برای تمامی اسمها با سطح پوشش 50/1% و صحتی معادل 99/1% برای صفتها (نسبت ها) ی نحوی و معنایی که از قوانین دست نوشته استفاده نموده بودند با سطح پوشش 73/6% بدست آورد (29).

4-5-1- ترکیب آماری شاخص های چندگانه

دو روش زیر شاخص های چندگانه را از نظر آماری ترکیب می نمایند تا معنای واژه بدست آید یعنی به جای تکیه بر روی یک نوع شاخص از چند نوع شاخص استفاده می نمایند .

1-5-4-1- مدل های تجزیه شدنی¹

بروس² و ویب³ مدلی را توصیف می کنند که از مشخصه های بافتی چندگانه استفاده نموده و نیازی به فرضیات آزمایش نشده ای در مورد شکل مدل ندارد، یعنی مشخصه هایی که ترکیب می شوند مستقل فرض نمیشوند (یا از نظر بافتی بسته به معنای واژه مستقل فرض نمی شوند) و برخی مدل های قبلی این فرضیه را ساخته اند . فرض اینکه تمام مشخصه ها در حقیقت به طریقی وابسته هستند بدان مفهوم است که داده های مربی بیشتری برای ساختن یک مدل خوب از یک واژه نیاز است. به جای آن، روش بروس و ویب مدلها را با تمام وابستگی های ممکن با استفاده از داده های مربی امتحان می نماید تا بفهمند کدام مدل انطباق بهتری دارد و کدام پارامترها را می توان متقابلاً مستقل فرض نمود. این مدل از چندین مشخصه که به دلیل ماهیت اطلاعی⁴ بودنشان انتخاب می شوند استفاده می کند. اگر مدلی که فرض می نماید ارزش مشخصه ای مستقل از معنای واژه است انطباق ضعیفی داشته باشد آن مشخصه ماهیت اطلاعی دارد. بهترین مدل آنها از مشخصه های زیر استفاده می نماید: مفرد / جمع، برجسب اجزاء کلام برای واژه مورد نظر و نیز دو واژه از هر دو طرف، و مشخصه های ویژه با هم آیی که برای هر واژه ای W که بخواهد ابهام زدایی شود متفاوت است.

¹ - Decomposable

² - Bruce

³ - Wiebe

⁴ - Informative

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

بروس و وب مدل خود را بر روی واژه interest امتحان کردند زیرا در مدل‌های قبلی این واژه یکی از سخت‌ترین واژه‌ها برای ابهام زدایی بود. مجموعه‌ای از داده‌ها شامل 2369 کاربرد واژه interest به عنوان اسم به طور دستی برچسب گذاری شد. 600 تا از اینها برای امتحان کنار گذاشته شد و بقیه برای آزمایش مورد استفاده قرار گرفتند و در کل، سیستم صحتی معادل 79% بدست آورد. بروس و وب مجموعه داده‌های خود را در اختیار عموم قرار دادند به طوریکه کارهای آینده بتواند مستقیماً با کار آنها مقایسه شود (30).

2-4-5-1- ابهام زدایی مبتنی بر نمونه

نگ¹ و لی² سیستمی را توصیف می‌نمایند که اطلاعات بدست آمده از مشخصه‌های چندگانه را نیز ترکیب می‌نماید. اما سیستم آنها از مدل‌هایی استفاده می‌نماید. این مدل‌ها مجموعه‌هایی از ارزشهای مشخصه‌ای هستند که نمونه‌هایی از واژه مورد نظر در داده‌های مربی را نمایش داده و سپس با نمونه واژه خواسته شده در متن با استفاده از معیار فاصله مقایسه می‌شود.

در نسخه اولیه این سیستم، واژه، معنایی که نزدیک‌ترین مدل دارد را به خود می‌گیرد و در نسخه اصلاح شده بعدی، K تا از نزدیک‌ترین نمونه‌ها پیدا شده و سپس متداول‌ترین معنی از این نمونه‌ها استفاده می‌شوند. مشخصه‌های مورد استفاده در این سیستم شبیه مشخصه‌های سیستم بروس و وب هستند، یعنی شامل اطلاعات مربوط به اجزاء کلام واژه مورد نظر و سه واژه از هر طرف، مشخصه صرفی (مفرد / جمع برای اسمها)، و کلید واژه‌هایی معادل با هم آیی‌های بروس و وب. نتایج بر روی داده interest که سیستم بروس و وب آن را امتحان کردند و نیز بر روی داده‌های پیچیده‌تری که شامل 121 اسم و 70 فعل بود عبارت بودند از صحتی معادل 89% در مورد داده اولی و 75/2% و 58/7% برای داده‌های دومی که به ترتیب از پیکره‌های زبانی WSJ و Brown استخراج شده بودند. این نتایج بدست آمده در مقایسه با الگوریتم دیگری به نام "نایویز3" که از همان مجموعه مشخصه‌ها استفاده می‌نمود بهتر به نظر می‌رسد (31).

¹ - Ng

² - Lee

³ - Naïve-Bayes

	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	تاریخ: 1388/04/7	ویرایش: 1/0	
	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ		

5-5-1- استفاده از پیکره یک زبانه زبان مقصد

سیستم‌های مختلفی بر این اساس پیشنهاد شده است که به چند نمونه از آنها اشاره کرده و آنها را مورد بررسی قرار می‌دهیم:

5-5-1-1- روش داگان^۱ و ایتای^۲

داگان و ایتای سیستمی را پیشنهاد می‌کنند که واژه‌های مبهم زبان مبدأ را با استفاده از داده‌های آماری مستخرج از پیکره زبانی زبان مقصد و یک واژه‌نامه دو زبانه (لیستی از واژه‌های معادل به زبان مبدأ و مقصد) ابهام زدایی می‌نماید. با استفاده از تجزیه‌گر زبان مبدأ الگوریتم آنها قادر به تشخیص ارتباطات نحوی بین واژه‌ها بوده و از فرهنگ دو زبانه برای منطبق کردن آن ارتباطات با امکانات مختلف (با این فرض که بیش از یک ترجمه ممکن برای برخی از واژه‌ها وجود دارد) در زبان مقصد استفاده می‌کند. سپس این امکانات مختلف در زبان مقصد بر طبق آمار بدست آمده از تعداد دفعات وقوع آنها در پیکره مربی انتخاب می‌شوند (32).

داگان و ایتای الگوریتم پیشنهادی خود را با استفاده از یک سیستم کامل مورد ارزیابی قرار ندادند زیرا برخی از عناصر مورد نیاز (مانند تجزیه‌گر و واژه‌نامه‌هایی برای زبان مقصد) در دسترس نبودند. تنها دو آزمایش انجام شد: یکی با استفاده از جملات عبری و دیگری با استفاده از جملات آلمانی. صحت بدست آمده برای سیستم، معادل 91% و 78% به ترتیب برای زبان عبری و زبان آلمانی بود.

سیستم داگان و ایتای تنها از واژه‌هایی که در قالب ارتباطات نحوی خاصی (مانند ارتباط بین فعل و مفعول و یا ارتباط بین فعل و فاعل) هستند استفاده می‌نماید. بنابراین اطلاعاتی را که می‌توان از دیگر واژه‌ها در بافت منطقه ای بدست آورد نادیده می‌گیرد. در برخی موارد که چندین ارتباط در رابطه با یک واژه وجود دارند ممکن است تصمیم‌گیری نهایی به دلیل انتخاب ارتباط نادرست یعنی ارتباطی که فراوانی‌های آن باعث انتخاب گزینه نادرست به جای گزینه درست می‌شود به سبب اینکه نمونه‌های بیشتری از آن در پیکره دیده شده است اشتباه باشد. می

¹ - Dagan

² - Itai

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

توان برای حل این مشکل ارتباطات نحوی را بار گذاری نمود بطوریکه ارتباطات نزدیک مانند ارتباط بین یک فعل و مفعول را جزء ارتباطات مهم تر قرار داد.

2-5-1- روش هیونگ لی^۱ و جنگ - سی پارک^۲ و گیل چنگ کیم^۳

در این روش آنها نیز از یک مدل آماری مبتنی برپیکره استفاده کردند و یک روش انتخاب واژه ای سه مرحله ای را پیشنهاد دادند. این مراحل عبارتند از: آشکار سازی معنی از کلمات مبدا، نگاشت معنی به کلمه، و انتخاب مناسب ترین جزء واژه ای زبان مقصد.

یک کلمه می تواند معانی بسیاری داشته باشد و هر کدام از معانی می تواند در کلمات زبان مقصد بسیاری نگاشت شود. دانشی و آگاهی لازم، برای هر مرحله از یک دیکشنری ماشین خوان و پیکره یک زبانه زبان مقصد اقتباس شده است. بوسیله عملکرد انتخاب واژه ای در سه مرحله و اقتباس آگاهی ضروری برای هر مرحله از منابع موجود، این سیستم کلمه مناسب برای ترجمه را با قابلیت تعمیم و نیرومندی بالا انتخاب می کند. در این روش اولین مرحله آشکار سازی کلمه از انتخاب واژه در جملات زبان مبدا با استفاده قوانین برای آشکار سازی معنی کلمه می باشد. قوانین زبان مبدا برای ابهام زدایی از معنی واژه^۴ از یک دیکشنری قابل خواندن ماشینی اقتباس می شوند و در دومین مرحله مفهوم انتخاب شده از کلمه زبان مبدا در مجموعه ای از کلمات زبان مقصد نگاشت می شود. هم چنین اطلاعات نگاشت شده از معنی کلمه زبان مقصد می تواند از یک دیکشنری قابل خواندن ماشینی اقتباس شود. یک کلمه زبان مقصد که بصورت مناسب در جمله زبان مقصد واقع شده است در

¹ - Hyun Ah Lee

² - Jong C. Park

³ - Gil Chang Kim

⁴ - Word Sense Disambiguation(WSD)

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

سومین مرحله انتخاب می‌گردد. برای انتخاب یک جزء واژه ای زبان مقصد، روابط اجزای واژه ای با یکدیگر در جملات زبان مقصد، با روابط صرفی و وزن‌های رابطه صرفی مورد استفاده قرار می‌گیرد. ترتیب کلمات در زبان مقصد و روابط نحوی شان از یک مجموعه پیکره‌ها یک زبانه زبان مقصد بصورت تجزیه و تحلیلی درختی فراهم شده اقتباس می‌شود. (33)

الف - آشکار سازی معنی کلمه مبدا و نگاشت کلمه

در اولین مرحله، معانی مختلف کلمه زبان مبدا با استفاده از دانش موجود در فرهنگ لغت مشخص می‌شود. در اینجا بر اساس معنی واژه در هر طبقه (یاروفسکی 1993)، از طبقه‌های کلمه مبدا، برای آشکار سازی معنی آن استفاده می‌شود و برای آشکار سازی دقیق‌تر، از اطلاعات رابطه ای صرفی بین کلمه مبدا و طبقه آن نیز استفاده می‌گردد. طبقات و رابطه صرفی شان در یک فرهنگ MT شرح داده شده است. به جدول زیر توجه نمایید.

جدول 1. قسمتی از دانش WSD برای کلمات Wear, Old

Table 1: Part of WSD knowledge for 'wear' and 'old'				
	relation	collocation	word sense	mapped target words
wear	object	coat, glasses	wear ¹ -1 (put on)	ip-ta, ssu-ta, sin-ta, cha-ta
		expression	wear ¹ -3 (express)	titi-ta, cis-ta
old	modifiee	man, book	old ¹ -2 (not young)	nai tun, nongyen-uy, nulk-un
		shoe, car	old ¹ -3 (not new)	nalk-un, hen

فرض کنید دوجمله مانند جملات زیر را به کامپیوتر داده تا آن را به زبان کره ای ترجمه نماید:

- (1) 'He wears an angry expression.'
- (2) 'She wears old shoes.'

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

در جمله اول معنی “Wear” پوشیدنی می باشد (یک عبارت خاصی درباره پوشاندن سطح) اما در جمله دوم معنی “Wear” پوشیدنی بر بدن می باشد. همچنین معنی “Old” در جمله دوم “غیر جدید” می باشد، اما در عبارت “Oldest student” کلمه “Old” به معنی “پیر” می باشد.

جدول فوق یک قطعه از دانش رفع ابهام از معنی کلمه موجود در فرهنگ لغت برای افعال Old و Wear را نشان می‌دهد.

منظور از Wear در جمله دوم “کفش ها” می باشد. بوسیله محاسبه کردن تشابه کفش و طبقه اش در جدول بالا می توان تعیین کرد که Wear در جمله دوم معنی $Wear^1 - 1$ را دارد. در اینجا از پیکره Wordnet برای محاسبه شباهت استفاده شده است. و دانش ابهام زدایی از معنای واژه به کمک دیکشنری، بطور اتوماتیک و و از یک دیکشنری ماشین خوان اقتباس شده است.

ب - نگاشت کلمه زبان مبدأ به مجموعه ای از کلمات زبان مقصد

در مرحله دوم مفهوم کلمه زبان مبدا انتخاب شده و به مجموعه ای از کلمات زبان مقصد که با آن معادل بوده ولی استفاده های متفاوت و تفاوت های جزئی دارند نگاشت می شود. ستون آخر جدول بالا اطلاعات نگاشتی از مفاهیم واژه مبدا در واژه های زبان مقصد را نشان می دهد. همچنین این اطلاعات نگاشتی از یک دیکشنری ماشین خوان اقتباس شده است. (33)

ج - انتخاب واژه با استفاده از اطلاعات آماری زبان مقصد

سومین مرحله، مرحله ای است که در بین واژه های زبان مقصد موردی را انتخاب می کند. یک کلمه مقصد که بصورت رایج در متن زبان مقصد مشابه رخ می دهد باید یک کلمه مقصد مناسب از مفهوم کلمه مبدا باشد. بافت زبان مقصد می تواند بوسیله جمع آوری کلمات ساخته شود. از اطلاعات آماری جمع آوری شده زبان مقصد برای انتخاب کلمه زبان مقصد مناسب استفاده می شود.

به جدول های 2 و 3 توجه نمایید:

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

جدول 2 برخی از کلمات را نشان می‌دهد که بصورت رایج همزمان با واژه مبهم در پیکره زبان مقصد رخ می

دهند.

جدول 2. کلماتی که بصورت رایج همزمان با واژه مبهم در پیکره مقصد رخ می‌دهند

Table 2: Sense of the object of wear¹-1 and its appropriate target word

	body-wear	head-wear	face-wear	hand-wear	foot-wear	arm-wear
wear ¹ -1	<i>sin-ta</i>	<i>ssu-ta</i>	<i>kki-ta</i>	<i>kki-ta</i>	<i>sin-ta</i>	<i>cha-ta</i>

جدول 3. بخشی از جمع تکرار کلمات مقصد Wear¹-1

Table 3: Part of cooccurrence statistic information

	"clothes"		"hat"	"glasses"	"glove"	"shoes"			"watch"
	<i>os</i>	<i>uypok</i>	<i>moca</i>	<i>ankyeng</i>	<i>cangkap</i>	<i>kwutwu</i>	<i>sinpal</i>	<i>sin</i>	<i>sikyey</i>
<i>sin-ta</i>	36	1	-	-	-	-	-	-	-
<i>ssu-ta</i>	1	-	8	11	-	-	-	-	-
<i>kki-ta</i>	-	-	-	2	2	-	-	-	-
<i>sin-ta</i>	-	-	-	-	-	2	5	1	-
<i>cha-ta</i>	-	-	-	-	-	-	-	-	1

و جدول 3 بخشی از جمع تکرار کلمات مقصد Wear¹-1 و دیگر کلمات موضوع را نشان می‌دهد. با کمک این جدول می‌توان فهمید که یک طبقه در برخی از رابطه‌های صرفی یک شاخص خوب برای انتخاب واژه‌ای خاص می‌باشد.

با استفاده از اطلاعات آماری زبان مقصد، تعداد هریک از ترکیب‌های نحوی در پیکره زبان مقصد مشخص می‌شود. (دایگان و ایتای 1994). دایگان و ایتای از ارتباطات صرفی و نحوی زبان مبداء برای ایجاد مجموعه چند تایی صرفی از رابطه صرفی زبان مقصد بکار می‌برند. اما در اینجا یک چند تایی صرفی با روابط صرفی زبان مقصد ساخته

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

می شود زیرا سومین مرحله انتخاب واژه بعد از انتقال ساختاری رخ می دهد. دو کلمه ای های زبان مقصد در مجموعه چند تایی صرفی عوامل مجموعه کلمات زبان مقصد هستند که در مرحله دوم نگاشت شده است. احتمال $p(a_i)$ را از انتخاب یک کلمه مقصد a_i به عنوان کلمه ای برای ترجمه ارزیابی می کنیم. بیایید فرض کنیم که مفهوم کلمه انتخابی a می باشد. ما می خواهیم یک کلمه را از a_1 و a_2 و ... a_k انتخاب کنیم که مجموعه ای از کلمات هدف مفهوم a می باشد. بیایند مجموعه از گره هایی¹ که به ساختار جمله زبان مقصد به گره معنی a متصل هستند با $\Theta(a)$ از یکدیگر تفکیک کنیم. اگر یک عنصر از مجموعه $\Theta(a)$ را بصورت گره (a, m, g) مشخص کنیم به این معنی که گره a می تواند به m کلمه زبان مقصد ترجمه شود، a_1 و a_2 و ... a_k ، و گره با رابطه صرفی g به گره با معنی a مرتبط شده است بنابراین این گره مجموعه ای از معانی a را نمایش میدهد و (g, a_i, a_1) یک مجموعه چند تایی صرفی روی معانی a_i می باشد. فرض کنید که $f(a_i, a_1, g)$ تعداد مجموعه چند تایی صرفی (g, a_i, a_1) در یک پیکره ها زبان مقصد می باشد. و $w(a, g)$ را بصورت وزن اهمیت رابطه نحوی g در انتخاب معنی کلمه a از زبان مقصد در نظر می گیریم. سپس احتمال $p^a(a_i)$ را که به مفهوم احتمال انتخاب a_i به عنوان معنی از کلمه مقصد a از فرمول زیر محاسبه شود:

$$n(a_i) = \sum_{(a, m, \gamma) \in \Theta(a)} \frac{\sum_{l=1}^m f(a_i, a_l, \gamma)}{m} \cdot w(a, \gamma) \quad (1)$$

$$p(a_i) = \frac{n(a_i)}{\sum_{j=1}^k n(a_j)} \quad (2)$$

رابطه 3-1- احتمال انتخاب a_i بعنوان معنی از کلمه مقصد a

$n(a_i)$ تعداد ظهور کلمه مقصد a_i در مجموعه های پیکره می باشد. با استفاده مجموعه $(a, m, \gamma) \in \Theta(a)$ ، معادله (1) برای جمع کلی تعداد همه گره هایی است که به گره با معنی a مرتبط شده

¹ - Nodes

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

است و بوسیله محاسبه تعداد کلی فرکانس‌ها با استفاده از رابطه نحوی (a, g, w) در (1)، ما می‌توانیم درجه اهمیت رابطه صرفی g را در انتخاب کلمه مقصدمفهوم a تخمین بزنیم.

جدول 4 احتمالهای بدست آمده را نشان می‌دهد و با استفاده فرکانس‌ها در جدول 4 و معادلات میتوان

$$n(\sin - ta) = \frac{1 + 5 + 2}{3} \times 1.0 + \frac{0}{1} \times 0.3 = 2.667$$

فرکانس را بدست می‌آوریم و احتمال

$p^{\wedge}(\sin - ta) = 0.748$ را بدست می‌آوریم سپس ما آستانه دینامیکی (پویایی) مشابه را برای $P1/p2$ از داگان و اتیای (1994) برای انتخاب بهترین کلمه برای ترجمه استفاده می‌کنیم. چون $p^{\wedge}(\sin - ta) / p^{\wedge}(ssu - ta)$ بزرگتر از مقدار آستانه‌ای در جمله 2 می‌باشد ما می‌توانیم کلمه مقصدمناسب را با استفاده از این روش انتخاب کنیم. (33)

3-5-5-1- روش موسوی و دلاورخلفی

در مدل آماری روش موسوی و دلاورخلفی به جای استفاده از مدل بیشترین احتمال¹ از اعداد تصادفی² و آزمون نکویی برازش³ استفاده نمودند. آنها معتقدند که استفاده از اعداد تصادفی برای انتخاب بهترین معادل زبان مقصد برای یک واژه مبهم از زبان مبدأ در سیستم ترجمه ماشینی نتایجی نزدیکتر به واقعیت در بر داشته و نسبت به هنگامی که انتخاب بر اساس محتمل‌ترین حالت صورت می‌گیرد دقیق‌تر و علمی‌تر است. فرض کنید جمله‌ای مانند جمله زیر را به کامپیوتر داده تا آن را ترجمه نماید:

She has a recently published article on biology.

از آنجا که واژه article واژه‌ای چند معنایی است در دیکشنری انگلیسی به فارسی چند معادل مختلف برای آن پیدا می‌کنیم از جمله آنها "مقاله، ماده و حرف تعریف". حتی اگر کامپیوتر بتواند با استفاده از روشهای مختلف برچسب گذاری اجزاء کلام طبقه‌بندی دستوری این واژه را نیز مشخص کند باز هم ابهام به قوت خود

¹- Most likelihood Model

²- random numbers

³- goodness of fit test

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

باقیست چراکه هر سه این معادل‌ها متعلق به طبقه‌ی دستوری یکسانی هستند، یعنی طبقه اسم. در اینجا است که ابهام زدایی معنی واژه (word sense disambiguation) معنا پیدا می‌کند.

اساس این روش بر این است که با توجه به جدول شماره 1 در مقاله نزدیک ترین اسم، صفت، یا فعل به اسم مورد نظربوسیله کامپیوتر تشخیص داده شده و معنی یا معناهای مختلف آن مشخص می‌شود. سپس فراوانیهای هم رخدادهای مختلف این دو واژه در پیکره یک زبانه زبان فارسی محاسبه می‌گردد. از آنجا که واژه published چندمعنایی نیست و واژه article دارای سه معنای مختلف می‌باشد تعداد هم رخدادهایی که باید محاسبه شوند $3 \times 1 = 3$ خواهد بود به شکل جدول 4 می‌باشد.

جدول 4. فراوانی هم رخدادها در پیکره

فراوانی در پیکره	هم رخدادها
X	۱- مقاله چاپ شده
Y	۲- ماده چاپ شده
Z	۳- حرف تعریف چاپ شده

با استفاده از فرمول:

$$P(tw_i) = \frac{f_i}{\sum_{j=1}^n f_i}$$

رابطه 4-1- احتمال وقوع هم رخدادها

و با استفاده از جدول توزیع تجربی و آزمونهای نکویی برازش توزیع احتمال هندسی X, Y و Z را بدست آورده و در نهایت برای هر گزینه یک احتمال وقوع مشخص می‌شود. در مدل بیشترین احتمال همیشه گزینه‌ای با بیشترین احتمال به عنوان گزینه درست انتخاب می‌شود. اما در مدل ما (اعداد تصادفی) این شانس به گزینه‌هایی

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

با احتمالات کمتر هم داده می‌شود تا با احتمال کمتری انتخاب شوند. در مثال بالا اگر احتمال وقوع 1 (مقاله چاپ شده) بیشتر از همه مثلاً 87 درصد باشد و احتمال وقوع 2 (ماده چاپ شده) مثلاً 7 درصد باشد، نباید در همه موارد 1 بعنوان گزینه درست انتخاب شود بلکه باید این امکان را به 2 نیز داد که در 9 درصد موارد انتخاب شود. این روش در مورد پیکره زبانی نسبتاً بزرگی در حوضه روانشناسی در زبانهای انگلیسی و فارسی به ترتیب به عنوان زبانهای مبدأ و مقصد آزمایش شد. الگوریتم این روش توانست مشکل ابهام 604 واژه از 764 واژه مبهم یک پیکره زبانی انگلیسی را با استفاده از متن معادل آن در زبان فارسی حل نماید که بدین ترتیب صحت این مدل به 79% رسید (12).

فصل دوم: کشف دانش

1-2- مقدمه ای بر کشف دانش

انسان همیشه برای الهام گرفتن به جهان زنده پیرامون خود نگریسته است. مشاهده ساختارهای اجتماعی آنان نتیجه مبادلات چند جانبه ای است که به ظاهر ساده اند. این مبادلات به کمک ارتباطهای شیمیایی صورت می‌گیرد که مورچه‌ها با شاخکهای خود این علائم را دریافت می‌کنند. فرومون ماده ای شیمیایی است که مورچه پس از عبور از خود برجا می‌گذارد. آزمایشی ساده این پدیده را ثابت می‌کند. پژوهشگران با قراردادن دو مسیر کوتاه و بلند در برابر مورچه‌ها برای رسیدن به غذا، پی بردند پس از مدتی همه مورچه‌ها از مسیر کوتاه تر عبور می‌کنند که علت آن تمرکز بیشتر فرومون در آن مسیر است. هر چه میزان تمرکز فرومون در مسیر کوتاه تر بیشتر باشد میزان تردد در آن بیشتر می‌شود. مورچه‌ها در طبیعت با این ماده می‌توانند کوتاه ترین مسیر را بین لانه خود و چندین منبع غذایی پیدا کنند. الگوریتم کلونی مورچه برای اولین بار توسط دوریگو¹ و همکارانش به عنوان یک راه حل چند عامله² برای مسائل مشکل بهینه سازی مثل فروشنده دوره گرد³ ارائه شد. عامل هوشمند⁴ موجودی

¹ Dorigo

² Multi Agent

³ TSP: Traveling Sales Person

⁴ Intelligent Agent

 دانشگاه گیلان	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

است که از طریق حسگرها قادر به درک پیرامون خود بوده و از طریق تاثیر گذارنده ها می تواند روی محیط تاثیر بگذارد. الگوریتم کلونی مورچه الهام گرفته شده از مطالعات و مشاهدات روی کلونی مورچه هاست. این مطالعات نشان داده که مورچه ها حشراتی اجتماعی هستند که در کلونی ها زندگی می کنند و رفتار آنها بیشتر در جهت بقا کلونی است تا در جهت بقا یک جزء از آن. یکی از مهمترین و جالبترین رفتار مورچه ها، رفتار آنها برای یافتن غذا است و بویژه چگونگی پیدا کردن کوتاهترین مسیر میان منابع غذایی و آشیانه. این نوع رفتار مورچه ها دارای نوعی هوشمندی توده ای است که اخیراً مورد توجه دانشمندان قرار گرفته است.

همانطور که می دانیم مسئله یافتن کوتاهترین مسیر، یک مسئله بهینه سازیست که گاه حل آن بسیار دشوار است و گاه نیز بسیار زمانبر. بعنوان مثال مسئله فروشنده دوره گرد (TSP). در این مسئله فروشنده دوره گرد باید از یک شهر شروع کرده، به شهرهای دیگر برود و سپس به شهر مبدا بازگردد بطوریکه از هر شهر فقط یکبار عبور کند و کوتاهترین مسیر را نیز طی کرده باشد. اگر تعداد این شهرها n باشد در حالت کلی این مسئله از مرتبه $(n-1)!$ است که برای فقط 21 شهر زمان واقعاً زیادی می برد:

$$\text{روز} = 1/7 * 10^{13} \text{ S} = 2/433 * 10^{16} \text{ ms} = 2/433 * 10^{18} \text{ ms} = 20!$$

با انجام یک الگوریتم برنامه سازی پویا برای این مسئله، زمان از مرتبه نمایی بدست می آید که آن هم مناسب نیست. البته الگوریتم های دیگری نیز ارائه شده ولی هیچ کدام کارایی مناسبی ندارند. ACO الگوریتم کامل و مناسبی برای حل مسئله TSP است.

2-2- بررسی موضوع و مسائل پیرامون آن

در این فصل ابتدا به منظور پیدا کردن درک یکسانی از مطالبی که در ادامه می آید به توضیح یک تعریف اصلی پرداخته می شود. سپس مسئله اصلی و هدف مورد نظر در این پروژه بیان می گردد.

 وزارت آموزش عالی و تحقیقات علمی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع عالی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن‌فارس — 2 — پ	

3-2- تعاریف اصلی

1-2-3-2- طبقه بندی چند هدفه و چند طبقه بودن¹

در صورتیکه نمونه‌هایی که برای آموزش بکار می‌روند، دارای n صفت و مقدار باشند که در واقع از این مقادیر برای تقسیم طبقه استفاده می‌شود ($m > 1$ و $n > 0$)، در این صورت فعالیت طبقه بندی روی این مجموعه‌ی آموزشی را طبقه بندی چند هدفه می‌گوییم. (34)

2-2-3-2- کشف دانش در پایگاه داده‌ها²

بعنوان استخراج اطلاعاتی مهم از میان داده‌ها، بطوریکه این اطلاعات قبلاً ناشناخته بوده و بطور بالقوه مفید باشند (17) تعریف گردیده است. تعریف متداول دیگری که برای کشف دانش در پایگاه داده آمده، عبارتست از: پردازش نه چندان ساده تشخیص الگوهایی نهایتاً قابل فهم در داده‌ها، بطوریکه معتبر، تازه و بالقوه مفید باشند. (35)

3-2-3-3- داده کاوی³

کاربرد الگوهایی خاص برای استخراج الگو از میان داده‌ها تعریف گردیده است. عبارت دیگر داده کاوی گامی از گامهای موجود در پردازش عمومی به منظور کشف دانش در پایگاه داده محسوب می‌گردد. (35) معادل‌های دیگری هم برای داده کاوی وجود دارد، منجمله: برداشت دانش⁴، کاوش دانش⁵، استخراج دانش⁶، و باستانشناسی داده⁷. (36)(35)

4-2-3-3- عامل⁸

¹ Multi Objective Classification

² Knowledge Discovery in Data Base(KDD)

³ Data-Mining

⁴ Information Harvesting

⁵ Knowledge mining

⁶ Knowledge extraction

⁷ Data Archaeology

⁸ Agent

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

موجودیتی خودگردان و محاسباتی که بتوان آنرا بعنوان ادراک کننده محیطش از طریق احساسگرها¹ و تأثیر گذار بر روی محیطش از طریق وسایل تأثیر دانست (37). یا سیستم کامپیوتری که در محیطی قرار گرفته و قادر به رفتار و فعالیتی خودگردان در آن محیط به منظور نیل به اهداف مورد طراحی اش باشد(38).

4-2- هدف پروژه

هدف از این پروژه بکارگیری الگوریتم یادگیری ماشین بر اساس الگوریتم بهینه سازی اجتماع مورچه ها² به منظور ابهام زدایی از واژه های چند معنایی است.

5-2- مروری بر کشف دانش در پایگاه داده ها

از آنجا که در این پایان نامه می خواهیم با استفاده از یکی از نمونه های موجود از الگوریتم های کشف دانش به ابهام زدایی از معنی واژه پردازیم ، لازم است ابتدا مروری بر کشف دانش در پایگاه داده ها داشته باشیم و بطور مختصر به بررسی دلائل و فناوری های موجود در این زمینه پردازیم.

1-5-2- دلائل توجه

دلائل توجه به کشف دانش در پایگاه داده ها را می توان در عوامل زیر خلاصه نمود: (36)(35)

- تعداد و اندازه پایگاه های داده در سازمانها با سرعت بسیار بالایی در حال رشد است. زمانی پتابایت و ترابایت ، حجمهای غیر قابل تصور محسوب می شدند در صورتیکه اکنون در حوزه های متفاوتی مانند فروش و حراج ، علوم زمین شناسی، مراقبت های بهداشتی و بیولوژی مولکولی (مانند پروژه ژنوم شناسی)و... به واقعیت تبدیل گردیده اند.
- سازمانها متوجه وجود دانش هایی با ارزش بالا در میان داده هایشان گردیده اند که در صورت کشف،مزیت رقابتی برایشان دربر خواهد داشت.
- بعضی از فناوری های مورد نیاز که توانایی عمل بر روی حجم های بالای داده ها را داشته باشند اخیراً به اندازه کافی رشد کرده اند.

¹ Sensors

² Ant Colony Optimization(ACO)

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

1-5-2- فناوری های کلیدی مورد استفاده

جستجوی دانش در پایگاه داده ، زمینه ای مربوط به رشته ها و فناوری های مختلف علمی است که هسته آن محل تلاقی یادگیری ماشین ، آمار و پایگاه داده هاست. (39) فناوری هایی که پیشرفت آنها کاربرد جستجوی دانش در پایگاه داده ها را امکان پذیر ساخته عبارتند از: (36)

- مدیریت پایگاه های داده
- هوش مصنوعی
- آمار
- یادگیری ماشین
- شناسایی الگو
- عاملهای هوشمند

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 رئیس‌جمهوری عالی اطلاع‌رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

فصل سوم: الگوریتم بهینه سازی

3-1- کشف دانش بوسیله الگوریتم بهینه سازی

برای کشف دانش گامهای متعددی بایستی صورت گیرد و در هر کدام چندین تکنیک برای اجرا موجود است. بر همین اساس، در گام مربوط به داده کاوی و شناسایی الگو از تکنیکهای متفاوتی می توان استفاده نمود که بعضی از آنها شامل طبقه بندی، خوشه بندی و مدل سازی مرتبط¹، تقلیل² و...مورد اشاره قرار گرفته است. هر کدام از این تکنیکها خود می توانند بعنوان وظیفه³ و مسئله ای جداگانه برای حل توسط الگوریتم کشف دانش دیده می شود (35). بنابراین اولین گام برای طراحی الگوریتم کشف دانش، مشخص ساختن وظیفه ایست که الگوریتم می خواهد آنرا اجرا نماید.

3-2- فعالیت کلاس بندی⁴

هدف از این فعالیت تخصیص هر مورد⁵ از اشیاء، رکوردها و یا نمونه ها⁶ به یک کلاس از میان مجموعه ای از طبقه های از پیش معین شده است. این فعالیت بر اساس مقادیری از بعضی صفات مشخصه (بنام صفات مشخصه پیشگو) برای آن مورد انجام می گیرد. دانشی که نهایتاً توسط کلاس بندی کشف می گردد و در ادامه این بحث مورد اشاره می باشند به شکل قوانین تصمیم بیان می گردد. شکل کلی این بیان دانش به صورت زیر می باشد:

¹ Dependence modeling

² Regression

³ task

⁴ Classification

⁵ Case

⁶ Instance

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن فارسی — 2 — پ	

IF <conditions> THEN <class>

مقدم قانون¹، شامل مجموعه‌ای از شرط‌هاست که معمولاً بوسیله‌ی عملگر منطقی عطف² بهم متصل گردیده اند. در ادامه این متن ما هر شرط را بعنوان یک عبارت³ مورد اشاره قرار می‌دهیم. بنابراین مقدم قانون، حاصل عطف منطقی عبارت‌هایی به شکل زیر است:

If term1 AND term2 AND...

هر عبارت متشکل از یک سه تایی <صفت مشخصه، عملگر، مقدار> است. مانند <جنسیت = مونث>. بخش پی‌آمد⁴، مشخص‌کننده کلاس پیش‌بینی شده برای مواردی است که صفات پیشگو، همه عبارت‌های مشخص شده در مقدم قانون را برآورده نمایند. از دید کشف دانش، این نحوه بازنمایی دانش این مزیت را دارد که توسط کاربر قابل درک است. البته این قابلیت درک تا زمانی است که قوانین کشف شده و تعداد عبارت‌های موجود در مقدم قوانین طولانی نباشند.

3-3- مزیت‌های الگوریتم بهینه سازی

همانطور که گفته شد «تبخیر شدن فرومون» و «احتمال-تصادف» به مورچه‌ها امکان پیدا کردن کوتاهترین مسیر را می‌دهند. این دو ویژگی باعث ایجاد انعطاف در حل هرگونه مسئله بهینه‌سازی می‌شوند. مثلاً در گراف شهرهای مسئله فروشنده دوره گرد، اگر یکی از یالها (یا گره‌ها) حذف شود الگوریتم این توانایی را دارد تا به سرعت مسیر بهینه را با توجه به شرایط جدید پیدا کند. به این ترتیب که اگر یال (یا گره‌ای) حذف شود دیگر لازم نیست که الگوریتم از ابتدا مسئله را حل کند بلکه از جایی که مسئله حل شده تا محل حذف یال (یا گره) هنوز بهترین مسیر را داریم، از این به بعد مورچه‌ها می‌توانند پس از مدت کوتاهی مسیر بهینه (کوتاهترین) را بیابند.

¹ If Part

² And

³ Term

⁴ Then Part

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

1-3-3- الگوریتم بهینه سازی

از آنجا که این الگوریتم الهام گرفته از فعالیت مورچه های واقعی در محیط طبیعی است ابتدا به بررسی رفتار مورچه های واقعی می پردازیم و بعد از آن به بیان الگوریتم می پردازیم.

2-3-3- حل مسئله

در یک کولونی از حشرات اجتماعی مانند مورچه ها ، زنبورها و موریانه ها هر حشره بطور معمول وظائف خاص خود را بصورت مجزا از سایر اعضاء کولونی انجام میدهد ، اما فعالیتهای انجام شده بوسیله ی هر حشره به شکلی با سایر اعضاء آن کولونی مرتبط می باشد بطوریکه در مجموع کولونی ایجاد شده توانایی و ظرفیت حل مسائل پیچیده را از طریق همکاری با هم دارا هستند (40).

فعالتهای مهم و حیاتی مانند انتخاب و جمع آوری مصالح مورد نیاز ، پیدا کردن و ذخیره سازی غذا که نیاز به طرح ریزی¹ پیچیده ای دارند ،همگی بوسیله ی کولونی های حشرات و بدون هیچ گونه نظارت خارجی و یا کنترل کننده مرکزی حل می گردند. این رفتار انتخابی که در حشرات و پرندگان اجتماعی دیده می شود تحت عنوان "هوش جمعی" شناخته می شود.

آنها هنگام انتخاب بین دو مسیر بصورت احتمالاتی مسیری را انتخاب می کنند که فرومون بیشتری داشته باشد یا عبارت دیگر مورچه های بیشتری قبلاً از آن عبور کرده باشند. حال دقت کنید که همین یک تمهید ساده چگونه منجر به پیدا کردن کوتاهترین مسیر خواهد شد .

فرض کنید مورچه ها روی مسیر AB در حرکت اند (در دو جهت مخالف) اگر در مسیر مورچه ها مانعی قرار دهیم مورچه ها دو راه برای انتخاب کردن دارند. اولین مورچه از A می آید و به C می رسد، در مسیر هیچ فرومونی نمی بیند بنابراین برای مسیر چپ و راست احتمال یکسان می دهد و بطور تصادفی و احتمالاتی مسیر CED را انتخاب می کند. اولین مورچه ای که مورچه اول را دنبال می کند زودتر از مورچه اولی که از مسیر CFD رفته به مقصد می رسد. مورچه ها در حال برگشت و به مرور زمان یک اثر بیشتر فرومون را روی CED حس می کنند و

¹ Planning

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

آنها بطور احتمالی و تصادفی (نه حتماً و قطعاً) انتخاب می کنند. در نهایت مسیر CED بعنوان مسیر کوتاهتر برگزیده می شود. در حقیقت چون طول مسیر CED کوتاهتر است زمان رفت و برگشت از آن هم کمتر می شود و در نتیجه مورچه های بیشتری نسبت به مسیر دیگر آنها طی خواهند کرد چون فرومون بیشتری در آن وجود دارد. در تصویر سمت چپ مورچه ها را در حال حرکت در مسیری مستقیم از لانه به سمت منبع غذایی مشاهده می شوند . تصویر سمت راست بیانگر چیزی است که در طبیعت بطور معمول اتفاق می افتد ، و آن قرارگیری یک شیء (مانع) در سر راه مورچه هاست . بمنظور دور زدن این مانع ، مورچه ها در ابتدا بصورت اتفاقی به چپ یا راست می پیچند (50%-50%).

نکته بسیار با اهمیت این است که هر چند احتمال انتخاب مسیر پر فرومون تر توسط مورچه ها بیشتر است ولی این کماکان، یک احتمال است و قطعیت ندارد. یعنی اگر مسیر CED پر فرومون تر از CFD باشد به هیچ عنوان نمی شود نتیجه گرفت که همه مورچه ها از مسیر CED عبور خواهند کرد بلکه تنها می توان گفت که مثلاً 90% مورچه ها از مسیر کوتاهتر عبور خواهند کرد. اگر فرض کنیم که بجای این احتمال قطعیت وجود می داشت، یعنی هر مورچه فقط و فقط مسیر پر فرومون تر را انتخاب میکرد آنگاه اساساً این روش ممکن نبود به جواب برسد. اگر تصادفاً اولین مورچه مسیر CFD (مسیر دورتر) را انتخاب می کرد و ردی از فرومون بر جای می گذاشت آنگاه همه مورچه ها بدنبال او حرکت می کردند و هیچ وقت کوتاهترین مسیر یافته نمی شد. بنابراین تصادف و احتمال نقش عمده ای در ACO بر عهده دارند.

نکته دیگر مسئله تبخیر شدن فرومون بر جای گذاشته شده است. بر فرض اگر مانع در مسیر AB برداشته شود و فرومون تبخیر نشود مورچه ها همان مسیر قبلی را طی خواهند کرد. ولی در حقیقت این طور نیست. تبخیر شدن فرومون و احتمال، به مورچه ها امکان پیدا کردن مسیر کوتاهتر جدید را می دهند. رفتار مورچه های واقعی را از میان حشرات دیگر و بخصوص توانایی آنها را در پیدا کردن کوتاهترین مسیر بین منبع غذایی و لانه شان¹ در محیط هایی که دائم در حال تغییرند و بدون داشتن دانش بصری ، مورد بررسی قرار می دهیم تا با الهام از آنها در محیط های کامپیوتری مسائل مورد نظر در زمینه ی کشف دانش حل گردد. (29)

¹ Nest

 دانشگاه گیلان	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی — 2 — پ	

در هنگامی که مسیری بین منبع غذایی و لانه مورچه موجود باشد و مورچه‌ها هم در اطراف یک شیء موجود حاضر باشند، آنها خیلی زود سعی در دور زدن آن مانع می‌کنند. در ابتدا هر مورچه راهی را به سمت چپ یا راست مانع با توزیع احتمال مساوی 50%-50% انتخاب می‌کند. از آنجا که همه مورچه‌ها با سرعتی نسبتاً یکسان حرکت می‌کنند و از خود مسیر فرومونی با نرخی ثابت از نظر میزان تولید فرومون باقی می‌گذارند، در نتیجه آن مورچه‌هایی که بصورت اتفاقی مسیر کوتاه‌تر منتهی به غذا را طی می‌کنند و در نتیجه سریعتر مسیرشان را نسبت به سایر مورچه‌ها می‌پیمایند و در نتیجه میزان فرومون هم در مسیر کوتاه‌تر سریعتر انباشته می‌گردد. چون مورچه‌ها تعقیب مسیرهایی با میزان فرومون بیشتر را ترجیح می‌دهند بنابراین بیشتر مورچه‌ها به مسیر کوتاه‌تر همگرا می‌شوند.

3-4- الگوریتم‌های بهینه‌سازی

یک الگوریتم بهینه‌سازی عموماً سیستمی بشکل عامل پایه¹ است که رفتار مورچه‌های طبیعی را شبیه‌سازی می‌کند و شامل مکانیزم‌هایی برای همکاری² و تطبیق³ می‌باشد. در مراجع (3) و (2) از این نوع سیستم بعنوان روشی فوق ابتکاری⁴ بمنظور حل مسائل ترکیبی⁵ بهینه‌سازی ارائه گردیده است. دو نوع کلی از الگوریتم بهینه‌سازی اجتماع مورچه‌ای وجود دارد که عبارتند از: سیستم کولونی مورچه‌ای⁶ و سیستم مورچه‌ای⁷، و در ادامه هر جا که سخن از ACO پیش می‌آید، منظور سیستم کولونی مورچه‌ای است. اثبات شده است که این روش فوق ابتکاری، هم مقاوم در برابر خطا⁸ و هم دارای تنوع⁹ است و به همین دلیل نیز بر روی محدوده‌ی وسیعی از مسائل بهینه‌سازی ترکیبی متفاوت، بطور موفق بکار رفته است (4). الگوریتم‌های بهینه‌سازی مبتنی بر کولونی مورچه، مبتنی بر ایده‌های اصلی زیر هستند:

¹ Agent base

² Cooperation

³ Adaptation

⁴ Meta-heuristic

⁵ Combinatorial

⁶ Ant Colony System (ACS)

⁷ Ant System

⁸ Robust

⁹ Versatile

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

- هر مسیر طی شده بوسیله ی یک مورچه برای یک مسئله ی فرضی با راه حل¹ کاندیدی برای آن مسئله در ارتباط است .
- وقتی که یک مورچه راهی را طی می کند ، میزان فرومون باقی گذاشته شده در آن مسیر با کیفیت راه حل² کاندید متناظر برای آن مسئله (هدف) خاص متناسب خواهد بود.
- هنگامیکه مورچه ای بین دو مسیر یا بیشتر ، قرار است مسیری را انتخاب می کند ، مسیر با مقدار فرومون بیشتر با احتمال بیشتری بوسیله ی آن مورچه مورد انتخاب می گیرد.
- در نتیجه ، نهایتاً مورچه ها به مسیر کوتاهتر همگرا خواهند شد ، که در این حالت مانند مورچه های واقعی برای بهینه بودن و یا نزدیک به بهینه¹ بودن راه حل یافته شده برای مسئله ی هدف ، امید زیادی وجود دارد .
- مراحل طراحی و تطبیق الگوریتم ACO برای هر مسئله به شرح زیر است: (20)
- نمایش درست و مناسبی از مسئله ی مورد نظر، و در صورت امکان نمایش به شکل گراف و یا ساختاری مشابه آن که بتواند بسادگی توسط مورچه ها پوشش داده شود ، مناسب خواهد بود.
- انتساب تابع هیوریستیک به هر کدام از انتخابهای مورچه در گام تولید یک راه حل
- روشی برای اطمینان پیدا کردن از ساخته شدن راه حل های معتبر
- در نظر گرفتن قانونی برای بروز رسانی مقدار فرومون (t)
- پایه گذاری راهی برای مقدار دهی اولیه ی فرومون
- قانون انتقال² مبتنی بر احتمالات بر اساس تابع هیوریستیک (h) و مقدار فرومون (t)، که بایستی مکرراً در ساخت قوانین بکار رود.

3-5- استفاده از الگوریتم در حل مسئله

برای پیاده سازی کامپیوتری رفتار مورچه های واقعی و تطبیق آن با مسئله ی کشف و استخراج دانش ، لازم است که هر مورچه را بعنوان یک عامل در نظر بگیریم . بنابراین هر مورچه ی مصنوعی بعنوان عاملی که رفتار مورچه ی واقعی را تقلید می کند ، تعریف می گردد (41)

¹ Near-Optimum

² Transition Rule

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

البته مورچه‌های مصنوعی و طبیعی دارای وجوه تمایز زیر هستند (42):

- مورچه‌های مصنوعی دارای حافظه هستند.
 - مورچه‌های مصنوعی کاملاً کور نیستند.
 - مورچه‌های مصنوعی در محیطی که زمان گسسته دارد، زندگی می‌کنند.
- با توجه به مطالب فوق مورچه‌های مصنوعی با مورچه‌های طبیعی دارای اشتراک‌هایی نیز می‌باشند که عبارتند از:
- مورچه‌های مصنوعی با بکارگیری یک روش مبتنی بر احتمال، مسیرهای با مقدار فرومون بیشتر را، به مسیرهای دیگر ترجیح می‌دهند.
 - مسیرهای کوتاه‌تر، دارای نرخ بالاتری از نظر میزان رشد فرومون می‌باشند.
 - مورچه‌ها از سیستمی غیر مستقیم برای برقراری ارتباط پایه‌ی میزان فرومون باقی گذاشته در هر کدام از مسیرها را استفاده می‌کنند.
- بنابراین یک سیستم کولونی مورچه‌ای (ACS) بصورت تکراری یک حلقه‌ی اصلی را انجام می‌دهد که در آن دو روال را فراخوانی می‌کند:

1. روال ساخت: که چگونگی ساخت و تغییر راه حل‌های مسئله‌ی در حال حل را مشخص می‌سازد.
2. روال بروز رسانی فرومونها در مسیرهای مختلف: که ارتباط غیر مستقیم بین مورچه‌ها را فراهم می‌نماید.

3-6- الگوریتم و جزئیات یادگیری

در الگوریتم ACO هر مورچه بصورت افزایشی¹ راه حلی را برای مسئله‌ی هدف می‌سازد و تغییر می‌دهد. در مورد کشف دانش، مسئله‌ی هدف، کشف قوانین کلاس بندی است. هر قانون کلاس بندی شکل زیر را دارد:

If <term1 AND term2 AND ...> THEN <Class>

هر عبارت، سه تایی <صفت مشخصه، عملگر، مقدار> است.

¹ Incremental

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

($\langle \text{value, operator, attribute} \rangle$) که value ی آن مقداری است که متعلق به دامنه ی صفت مشخصه .
المان عملگر در هر سه تایی یکی از عملگرهای نسب¹ می تواند باشد.
نسخه ای از الگوریتم کشف دانش بوسیله ی کولونی مورچه ها که در ادامه مورد بررسی قرار می گیرد برای
کار با صفات تطبیق داده شده است و بنابراین عضو عملگر در سه تایی فوق همواره "=" خواهد بود.

3-7- سیستم یادگیری

در این قسمت سیستم مورچه معرفی شده و از مسئله ی فروشنده دوره گرد بعنوان معیار استفاده می گردد. در
مسئله فروشنده دوره گرد ، یک فروشنده سفر خود را از یک شهر آغاز کرده و پس از یک سفر کامل دوباره به
شهر خود باز می گردد و از هر شهر فقط یکبار عبور می کند و در ضمن باید از همه ی شهرها هم عبور کند .
هدف یافتن کوتاهترین مسیر برای این سفر می باشد .

d تعداد شهرها می باشد و dij فاصله اقلیدسی بین دو شهر i و j برابر است با:

$$dij = [(x_i - x_j)^2 + (y_i - y_j)^2]^{0.5}$$

رابطه 3-1 dij فاصله اقلیدسی بین دو شهر i و j

¹ Relational operator

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیر پروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک-متن-فارس — 2 — پ	

یک نمونه از TSP بوسیله ی یک گراف (N, E) که N مجموعه ی شهرها و E مجموعه ی یالها (فاصله های) میان شهرها می باشد. الگوریتم مورچه برای حل مسئله ی فروشنده ی دوره گرد بکار رفته است.

b_i , $i = 1, 2, 3, \dots, d$ را تعداد جمعیت (مورچه ها) در شهر i - ام در زمان t و $\sum_{i=1}^d b_i(t)$ کل جمعیت

مورچه ها می باشد. هر مورچه یک نماینده ی ساده با ویژگیهای زیر است:

1. آن یک شهر را برای رفتن انتخاب می کند که تابعی از فاصله شهر و مقدار اثر فرومون موجود در آن مسیر می باشد.

2. برای وادار کردن مورچه ها جهت انجام سفرهای منطقی، سفر به شهرهایی که یکبار از آن عبور کرده اند ممنوع می باشد.

3. هنگامی که یک مورچه یک سفر کامل را انجام می دهد مقداری فرومون بر روی هر مسیر i و j باقی می گذارد.

t_{ij} شدت اثر فرومون در مسیرهای (i, j) در زمان t می باشد. هر مورچه در زمان t شهر بعدی را انتخاب می کند که در زمان $t+1$ در آن خواهد بود. بنابراین اگر یک تکرار از الگوریتم سیستم مورچه m حرکت انجام شده بوسیله m مورچه در فاصله زمانی $(t, t+1)$ در نظر گرفته شود، بعد از d تکرار الگوریتم (یک سیکل) هر مورچه یک سفر کامل انجام داده است. در این نقطه شدت اثر فرومون بر اساس رابطه ی $t_{ij}(t+1) = r \cdot t_{ij}(t) + \Delta t_{ij}$ محاسبه می شود و r یک ضریب است بطوریکه (r^{-1}) میزان تبخیر فرومون را در فاصله زمانی $(t, t+1)$ را نشان

می دهد و $\Delta t_{ij} = \sum_{k=1}^m \Delta t_{ij}^k$ که مقدار اثر فرومون در هر واحد طول می باشد که بر روی مسیر (i, j) بوسیله ی k - امین مورچه در فاصله ی $(t, t+1)$ بر جای گذاشته می شود. قابل ذکر است که سه نوع الگوریتم مورچه وجود دارد:

1. Ant densit
2. Ant quantityy
3. Ant cycle

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در دو الگوریتم اول مقدار فرومون در هر تکرار تعدیل و بروز رسانی می شود اما در الگوریتم سوم بعد از یک سیکل عمل تعدیل صورت می گیرد .

در الگوریتم اول داریم :

$$\Delta t_{ij}^k(t, t+1) = \begin{cases} Q_1 & \text{اگر } k \text{ اُمین مورچه در فاصله زمانی } (t, t+1) \text{ از } i \text{ به } j \text{ می رود} \\ 0 & \text{در غیر این صورت} \end{cases}$$

که Q_1 یک مقدار ثابت است که در واحد طول در مسیر (i, j) بر جای گذاشته می شود.

در الگوریتم دوم :

$$\Delta t_{ij}^k(t, t+1) = \begin{cases} Q_2 / d_{ij} & \text{اگر } k \text{ اُمین مورچه در فاصله زمانی } (t, t+1) \text{ از } i \text{ به } j \text{ می رود} \\ 0 & \text{در غیر این صورت} \end{cases}$$

که Q_2 یک مقدار ثابت و d_{ij} فاصله ی بین شهر i و j میباشد.

در الگوریتم سوم :

$$\Delta t_{ij}^k(t, t+1) = \begin{cases} Q_3 / L^k & \text{اگر } k \text{ اُمین مورچه از مسیر های } (i, j) \text{ در سفرش استفاده کند} \\ 0 & \text{در غیر این صورت} \end{cases}$$

Q_3 مقدار ثابت و d_{ij} طول سفر k -اُمین مورچه می باشد.

 وزارت آموزش عالی و تحقیقات علمی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک-متن فارسی 2 — پ	

ضریب باید دارای ارزش کوچکتر از یک باشد تا از انباشتگی نامحدود فرومون جلوگیری کند. (46)

فصل چهارم: کشف دانش چند طبقه

4-1- کشف دانش چند طبقه مبتنی بر بهینه سازی

قبل از پرداختن به کارهای انجام شده توسط این بخش از پروژه، لازم است که درباره ی دانش چند طبقه و نحوه ی نمایش آن وسایر مطالب مربوط به مقدمات کار به زبانی مشترک دست یابیم. از اینرو ابتدا بخشی جداگانه برای معرفی نحوه ی نمایش دانش چند طبقه آورده می شود.

4-2- دانش چند طبقه و نحوه نمایش آن

همانطور که می دانیم حرف از کشف دانش چند طبقه، استخراج دانش با طبقه ای چند بعدی است. شکل 5 نمایش دهنده ی چند طبقه بودن یک قانون است.

```

IF term11 AND term12 .....AND term1M THEN Class 1
IF term21 AND term22 .....AND term2M THEN Class 2
.....
IF termi1 AND termi2 .....AND termiM THEN Class i
.....
IF termk1 AND termk2 .....AND termkM THEN Class k

```

شکل 5. تعبیر یک قانون چند طبقه

در این شکل هر قانون چند طبقه به معنی آرایه ای از قوانین تک طبقه فرض شده است. مقدار k معادل تعداد طبقه های موجود در مجموعه ی آموزشی است. مقدار M برابر تعداد صفات مشخصه برای هر نمونه آموزشی است. به عبارت ساده یک قانون چند طبقه معادل آرایه ای از قوانین یک طبقه ای فرض شده است. تعداد قوانین یک طبقه ای در هر قانون چند طبقه، معادل تعداد طبقه های موجود در مجموعه ی آموزشی است.

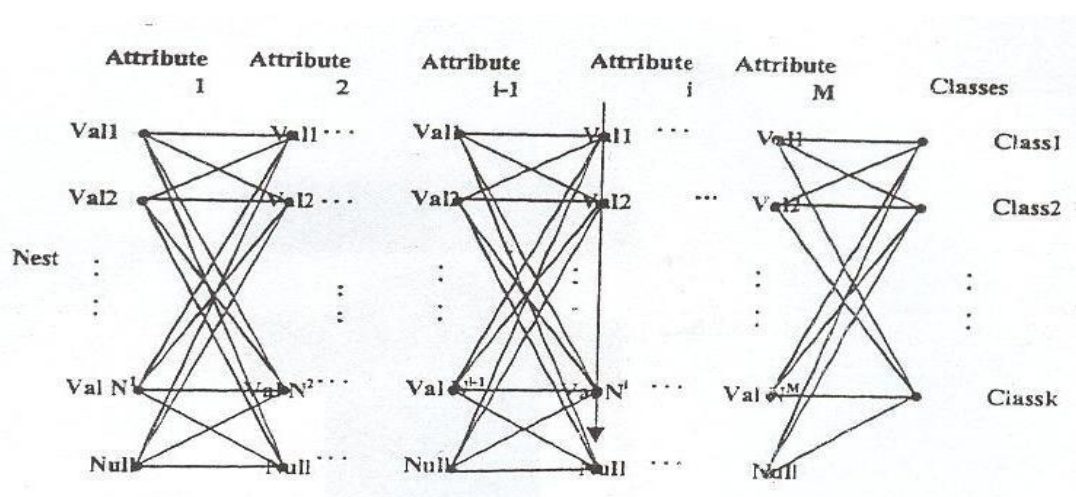
 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع ملی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن‌فارس — 2 — پ	

این ظاهر شدن قوانین از همه ی طبقه ها باعث می شود از داشتن دانش مربوط به نمونه هایی که تعدادشان در مجموعه ی آموزشی کم است ؛ اطمینان یابیم. توجه به این نکته لازم است که قوانین یک طبقه موجود در قانون چند طبقه ای نمی توانند بخشهای مقدم معادل هم داشته باشند.

3-4- تحلیل هدف از نقطه نظر الگوریتم بهینه سازی

هدف کشف دانش چند طبقه (پیدا کردن معنی مناسب واژه از بین معانی مختلف واژه) با استفاده از این الگوریتم است. توجه به گامهای استاندارد در ساخت هر الگوریتم جدید مبتنی بر ACO اهمیت زیادی داشته و بر سرعت طراحی می افزاید.

ساده ترین ساختار برای پیمایش توسط مورچه ها ، ساختار گراف است . از اینرو ابتدا مسئله را با این ساختار نمایش می دهیم . شکل 3 در زیر نمایش دهنده فضای مسئله به شکل گراف است .



شکل 6. نمایش فضای مسئله با استفاده از گراف

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در شکل 6 هر گره گراف معادل یک عبارت¹ در نظر گرفته شده و هر بخش عمودی از گره‌ها (ستون) معادل مجموعه مقادیر متعلق به دامنه‌ی یک صفت مشخصه است. مقدار Null در دامنه‌ی هر صفت مشخصه نشان دهنده‌ی انتخاب نشدن آن صفت مشخصه است. بنابراین هر مسیر از لانه به یک کلاس، معادل یک قانون یک طبقه‌ای برای انتخاب آن طبقه می‌تواند تلقی شود.

توجه به این نکته لازم است که برای جواب قابل قبول قانون چند طبقه، وجود فقط یک مسیر از لانه به کلاس مورد نظر کافی نیست و بایستی مسیرهای به تمام طبقه‌ها را یافت؛ یعنی K مسیر در هر بار. با توضیحات فوق می‌توان مسئله‌ی کشف دانش چند طبقه را معادل کشف K قانون یک طبقه‌ای با فرض وجود محدودیتهای زیر دانست:

(k) قانون یک طبقه‌ای کشف شده بایستی تمام (k) کلاس مجموعه آموزشی را داشته باشند، یعنی (k) قانون یک طبقه‌ای کشف شده بایستی کلاس تکراری در بخش نتیجه قوانین نداشته باشند. بخش مقدم (k) قانون یک طبقه‌ای کشف شده بایستی معادل هم باشند.

از نظر گراف، راه حل مسئله کشف دانش چند طبقه، بصورت داشتن k مسیر از لانه به k طبقه (هدف) به صورتی که این مسیرها تکراری نباشند و اهداف تکراری هم نداشته باشند، تعریف می‌گردد. حال که توانستیم مسئله را بصورت یک طبقه‌ای در نظر بگیریم بنابر این می‌توان از روش مربوط به کشف دانش یک طبقه‌ای کمک گرفت. البته اینکار بایستی با توجه به محدودیتهای اعمالی فوق صورت گیرد. اما بایستی به این نکته توجه داشت که برای هر بار بدست آوردن دانش یک طبقه‌ای در فضای ترسیمی شکل 2، گراف بایستی پیمایش گردد. پس کشف دانش چند طبقه معادل پیمایش‌های متوالی و یا موازی به تعداد k بار خواهد بود.

اگر فضای را بصورت k لایه بر روی هم در نظر بگیریم یعنی k- تا از شکل 5 بر روی هم، می‌توانیم با تغییراتی اندک از روش گفته شده قبلی برای کشف دانش یک طبقه‌ای با استفاده از این الگوریتم در کشف دانش چند طبقه‌ای یاری بگیریم.

¹ Term

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 رئیس‌رای عاين اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

البته همانطوری که گفته شد می توان صرفنظر از چند لایه ای شدن ، بر اساس شکل 5 و بصورت متوالی این مسئله را حل کرد. در این حالت مسئله را بصورت راه حل های جداگانه در نظر می گیریم که در هر دور از حل ، یک قانون یک طبقه ای کشف می گردد. ولی این نوع تحلیل با کشف کشف قانون چند طبقه بر اساس توضیح گفته شده قبلی برای چند طبقه ای بودن ، چندان مناسبی ندارد زیرا برای چند طبقه ای بودن لازمه ی عمل کشف قانون چند طبقه ، ظاهر شدن قانونهای یک طبقه ای از همه ی طبقات بصورتی یکجا است . لذا در این پروژه فضای مسئله به شکل k لایه ای در نظر گرفته می شود. k لایه ای بودن از نظر تداخل نداشتن فرومونها برای هر طبقه با طبقه ی دیگر لازم بنظر می رسد در غیر این صورت مسئله را باید در k بار و بصورت جداگانه حل نمود و برای هر بار حل ، مسیرها را فرومون دهی اولیه نمود.

سایر گامهای طراحی عبارتند از: انتساب تابع هیورستیک ، پیدا کردن قانون مناسبی برای بروز رسانی فرومون (t) و...

بدلیل آنکه می خواهیم از روش ACO در کشف دانش یک طبقه ای استفاده نماییم لذا بایستی نمایش فضای مسئله را بصورت یک طبقه ای تبدیل نماییم . این تبدیل از دو جنبه حائز اهمیت است :

1. از حیث انطباق بیشتر مسئله
 2. درک بهتر مسئله کشف دانش چند طبقه ، زیرا ترسیم و تصور گراف شکل 5 بصورت k لایه بر روی هم چندان هم ساده نیست.
- برای انجام دادن تبدیل ، بایستی هدف از چند لایه ای شدن را تحلیل نمود . تبدیل شکل بالا به حالتی k -لایه ای باعث می شود که بتوان مورچه هایی را تصور نمود که در لایه های جداگانه ، بطور موازی برای رسیدن به اهداف مختلف (k -هدف) فعالیت داشته باشند . این فعالیت بر روی لایه های جداگانه باعث می شود که هر گروه از مورچه در هر لایه ، را بصورت عواملی متخصص در رسیدن به هدف خاص آن لایه در نظر گرفت.
- پس برای تبدیل فضای مسئله یعنی تبدیل گراف k -لایه ای شکل بالا لازم است که تعریف جدیدی از مورچه داشته باشیم که با دید قبلی ما از مورچه کمی متفاوت است . بعبارت دیگر لازم است مورچه ها را متخصص در نظر بگیریم تا تبدیل فضای مسئله ، بطور کامل درک شود.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

1-3-4- گامهای دیگر طراحی

بدلیل آنکه مسئله ی چند طبقه ای به حل مسئله های یک طبقه ای همزمان تبدیل شد، بنابراین بقیه ی گامهای طراحی مانند تابع هیوریستیک ، قانونی برای بروز رسانی فرومون (t) و ... همان هایی خواهند بود که در مسائل یک طبقه ای مطرح میشود که در ادامه و در هنگام توضیح الگوریتم به آن می پردازیم. این تبدیل از دو لحاظ دارای اهمیت است:

1. از حیث انطباق بیشتر مسئله جاری با مسئله حل شده قبلی
 2. درک بهتر مسئله کشف دانش چند طبقه ، زیرا ترسیم و تصور گراف شکل 5 بصورت k - لایه بر روی هم چندان ساده نیست.
- برای انجام دادن تبدیل، بایستی هدف از چند لایه ای شدن شکل 5 را تحلیل نمود. تبدیل شکل 2 به حالتی k - لایه ای باعث می شود که بتوان مورچه هایی را تصور نمود که در لایه های جداگانه، بطور موازی برای رسیدن به اهداف مختلف (k -هدف) فعالیت داشته باشند. این فعالیت بر روی لایه های جداگانه باعث می شود که هر گروه از مورچه در هر لایه را بصورت عواملی متخصص در رسیدن به هدف خاص آن لایه در نظر گرفت.
- پس برای تبدیل فضای مسئله یعنی تبدیل گراف K - لایه ای شکل 5 به شکلی معادل شکل 7 که در الگوریتمهای دیگر استفاده شد، لازم است که تعریف جدیدی از مورچه را داشته باشیم که با دید قبلی که از مورچه داشته ایم ، کمی متفاوت است. بعبارت دیگر لازم است ، مورچه ها را متخصص در نظر بگیریم تا تبدیل فضای مسئله ، کاملاً درک گردد.

4-4- سلسله مراتب در الگوریتم

فرض تخصصی شدن کار مورچه های مصنوعی با فعالیت مورچه ها در طبیعت سازگار است. زیرا همانطور که می دانیم مورچه ها در طبیعت دارای سلسله مراتب و تخصصهای خاص خود هستند. بعنوان مثال مورچه ها را می توان به مورچه های کارگر، مورچه های سرباز و ملکه تقسیم بندی نمود. که البته هر کدام از این دسته ها دارای زمینه ی تخصصی خاص خود هستند. (43) (44) (45). هرچه میزان پیچیدگی و تعداد مورچه های کولونی بیشتر

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیر پروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن‌فارس — 2 — پ	

شود ، تخصصهای مورچه ها هم زیادتر می گردد. تعداد مورچه ها می تواند بین حداقل 30 تا میلیون مورچه را در برگیرد.(44) اینکه آیا نحوه ی ارتباط برای مورچه ها از یک تخصص با مورچه های دیگری از تخصص های دیگر بصورتی متفاوت یا یکسان است ،مطلبی است که به جانور شناسی مربوط است و ما فعلاً با این مسئله کاری نداریم .

مشخصات مورچه های متخصص مصنوعی بصورت زیر است :

1. مورچه های مصنوعی از یک تخصص با مورچه های دیگر از تخصص های دیگر هیچگونه ارتباطی ندارند.
2. روش ارتباطی مورچه ها در یک تخصص بر اساس فرومون و بر اساس قواعد گفته شده می باشد.
3. تخصص مورچه های مصنوعی امری ذاتی است و از ابتدای ایجاد آنها با آنها بوجود می آید . بعبارت دیگر نمی توان تخصص یک مورچه را در خلال زندگی آن تغییر داد.

1-4-4- پی آمد ناشی از تخصصی شدن کار مورچه ها

با فرض تخصصی شدن کار مورچه ها ، هر گروه تخصصی ، بر روی لایه ای جداگانه و مخصوص به تخصصشان کار می کنند. امکان تبادل اطلاعات بین مورچه های متعلق به تخصصهای متفاوت متصور نیست ، زیرا اطلاعات راهنما در ساخت قانون یک کلاس برای ساخت قانون کلاس دیگر مفید نیست . بنابراین جدا نمودن لایه ها در اصل سعی بر جداسازی اطلاعات مبادله شده بین مورچه های هر تخصص را دارد بدون آنکه بر روی هم تأثیر منفی داشته باشند.

پی آمد دیگری که بایستی به آن توجه نمود ثابت ماندن تخصص هر مورچه ی متخصص در دوره زندگی آنست .

2-4-4- تغییر فضای مسئله

همانطور که گفته شد جدا نمودن لایه ها، در اصل سعی بر جداسازی اطلاعات مبادله شده بین مورچه های هر تخصص را دارد . اطلاعات مبادله شده میان مورچه ها در محاسبه ی احتمال انتسابی به هر عبارت متبلور می شود .

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیر پروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن فارسی — 2 — پ	

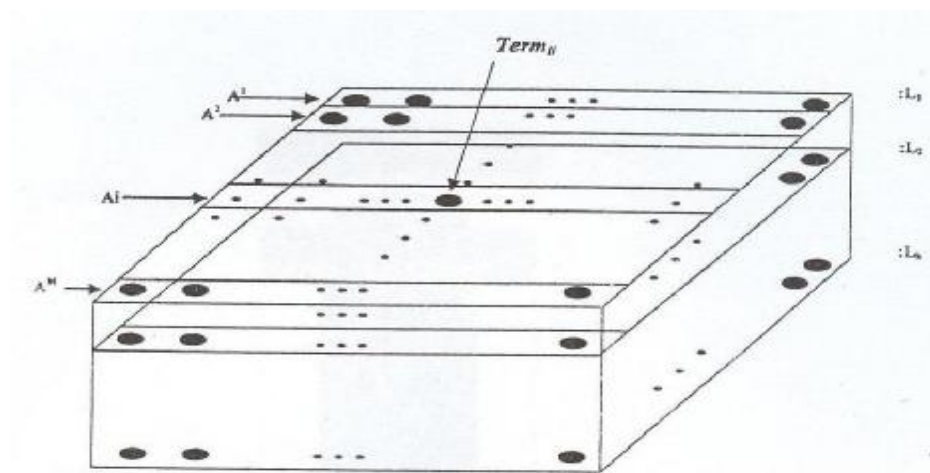
هر احتمال انتسابی حاوی اطلاعاتی از فرومون و تابع هیوریستیک بصورت توأم است. بنابراین بایستی جداسازی و لایه ای نمودن بر روی احتماله‌های عبارات صورت گیرد.

اکنون برای چند لایه ای نمودن شکل 2 کافیهست k تا جدول شکل 4 را تصور نمود.

در این حالت، در هر حالت، در هر لایه، هر عبارت دارای احتمال انتسابی متفاوتی با لایه های دیگر است. این تفاوت ناشی از تفاوت اطلاعات مبادله شده در هر لایه است.

شکل زیرنمایش جدول خصیصه و مقدار برای گراف شکل بالا را در حالت k - لایه ای نمایش می دهد.

بعبارت دیگر فضای مسئله کشف دانش k - طبقه بصورت k - جدول خصیصه و مقدار نمایش داده شده است، تا بتواند حالت حجمی بیابد.



شکل 7. جدول خصیصه و مقدار برای گراف k - لایه ای شکل 3

توضیح آرگومانهای شکل 7 بصورت زیر است:

	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
	تاریخ: 1388/04/7	ویرایش: 1/0	
	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ		

A_i نمایش دهنده ی صفت مشخصه ی I - A_i است .
 $(term_{ij})$ که معنی آن عبارتی با شرط $A_i = V_{ij}$ است ، که A_i صفت مشخصه i - A_i است و V_{ij} مقدار j - A_i از دامنه A_i .

مقدار L_k نشان دهنده لایه ی k - A_i از جدول است .
 مقدار k برابر تعداد طبقه های موجود در مجموعه ی آموزشی است.
 بایستی به این نکته توجه داشت که هر لایه از شکل 7 فضای کاری یک دسته از مورچه های متخصص را تشکیل می دهد . یعنی مورچه های متخصص در بدست آوردن کلاس i - A_i در لایه i - A_i A_i مشغول می شوند. همزمان با آنها سایر مورچه ها از سایر تخصصها در حال کار بر روی لایه های متناظر با تخصصشان می باشند . این همزمانی فعالیت مورچه ها از مهارت های متفاوت برای کسب اهدافی متفاوت در لایه های مختلف به معنی توان پردازشی موازی بالاتر برای الگوریتم کشف دانش چند طبقه مبتنی بر کولونی مورچه هاست.

5-4- بررسی چگونگی انتخاب عبارات و نمایش مسیر های مورچه ها

همانطور یکه در الگوریتم زیر خواهیم دید ، مورچه ها در حلقه ی Repeat-Until فعالیت می کنند و اصلی ترین فعالیت آنها ، ساخت قانون است. این فعالیت به صورت انتخابهای مکرر عبارت ها از میان عبارتهای موجود ی که قبلاً انتخاب نشده اند و اضافه کردن عبارت انتخابی به قانون در حال شکل گرفتن بیان می شود.
 شکل زیر نمایش دهنده ی عبارت های محتمل برای انتخاب توسط مورچه هاست ، عبارت دیگر فضای کلیه ی حالات ممکن برای انتخاب توسط مورچه ها را نمایش می دهد:

Attribute1: value1(P_{11}) , value2(P_{12}) , ...,value $N^1(P_{1N}^1)$
 Attribute2: value1(P_{21}) , value2(P_{22}) , ...,value $N^2(P_{2N}^2)$

 AttributeM: value1(P_{M1}) , value2(P_{M2}) , ...,value $N^M(P_{MN}^M)$

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

شکل 5. نمایش دهنده عبارت های محتمل

پارامترهای شکل 5 در زیر تعریف گردیده است:

Attribute j: صفت مشخصه j-ام در قانون.

N^i : تعداد مقادیر دامنه برای صفت مشخصه i-ام (Attribute j).

Value k: مقدار k-ام از دامنه ی صفت مورد نظر.

P_{ab} : احتمال انتخاب شدن مقدار b-ام از صفت مشخصه a-ام. بعبارت دیگر احتمال انتخاب عبارت ($term_{ab}$) که به شکل $A_a = V_{ab}$ نمایش داده می شود.

توجه به این نکته لازم است که تعداد حالات محتمل برای انتخاب شدن بوسیله ی یک مورچه در ابتدای کار ساخت قوانین ، عددی معادل $(N^1 + 1) * (N^2 + 1) * ... * (N^M + 1)$ (حاصل ضرب تعداد مقادیر دامنه بعلاوه یک برای هر صفت مشخصه) است.

هر مورچه می تواند در هر بار یک مقدار از مقادیر مربوط به یک صفت مشخصه را بردارد، بعد از آن دیگر حق برداشتن هیچ مقداری از آن صفت را ندارد. انتخاب شدن عبارت ها و اضافه شدن آنها به قانون در حال ساخت در مراجع (41) و (42) بعنوان ایجاد مسیر جزئی¹ توسط مورچه ها در نظر گرفته شده است.

4-6- نحوه کار الگوریتم چرخ رولت در انتخاب عبارت ها

حال به بررسی مکانیزم انتخاب عبارت ها در الگوریتم می پردازیم . در روش انتخابی عبارات بر اساس چرخ رولت ، هر عبارت بر اساس احتمال انتسابی به آن ، فضایی از محیط چرخ را به خود اختصاص می دهد . بعد از تقسیم بندی محیط چرخ برای آن عددی تصادفی تولید می کنیم (که معادل چرخاندن چرخ است). اگر عدد ایجاد شده (شاخص) با هر یک از بخشهای مختلف بوجود آمده بر روی محیط تلاقی کرد ، آنرا بعنوان عبارت مورد نظر انتخاب می کنیم .

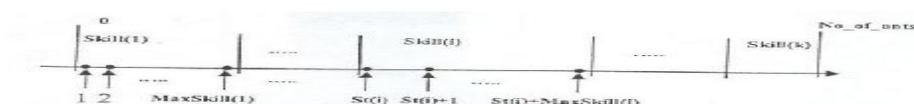
¹ Partial path

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

مزیت این روش در سادگی پیاده سازی و رعایت عدالت است. یعنی این روش به عبارتهای مختلف بر اساس میزان احتمال انتسابی به آنها، شانس انتخاب شدن می دهد.

7-4- الگوریتم کشف دانش چند طبقه مبتنی بر ACO

با توجه به مطالبی که در رابطه با دانش چند طبقه و معادل بودن آن با کشف چندین قانون یک طبقه ای گفته شد الگوریتم زیر کشف دانش چند طبقه مبتنی بر الگوریتم بهینه سازی کولونی مورچه ها را نمایش می دهد. قبل از پرداختن به الگوریتم لازم است نحوه ی دسترسی به هر مورچه متخصص را مشخص نماییم. هر مورچه در کولونی با اندیس منحصر بفرد و مربوط به خود شناسایی می گردد. هر مهارت از مورچه ها از شماره اندیس خاصی در کولونی شروع می شود و به شماره اندیسی در کولونی ختم می گردد. شماره مورچه های متخصص از یک مهارت بصورت متوالی است. به شکل 8 توجه نمایید:



شکل 8. محیط چرخ رولت برابر عدد یک است. احتمالات مربوط به عبارت ها، محیط چرخ رولت را به بخشهایی تقسیم نموده است. عددی اتفاقی یکی از بخشها (segment) روی این محیط را انتخاب کرده است. این عدد اتفاقی مانند یک شاخص بر روی محیط چرخ رولت عمل می کند.

در این شکل k تخصص و مهارت جداگانه وجود دارد. مقدار k مساوی است با تعداد طبقه های موجود در مجموعه ی آموزشی. تخصص i -ام از این کولونی در میانه ی جمعیت در شکل 4 مشخص شده است. $St(i)$ اندیس اولین مورچه در مهارت i -ام است. $MaxSkill(i)$ حداکثر مورچه های ممکن متصور برای مهارت i -ام را نمایش می دهد. اندیس آخرین مورچه از تخصص i -ام برابر با $St(i) + MaxSkill(i)$ خواهد بود.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:			 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه:			
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در ادامه الگوریتم کشف دانش چند طبقه ی مبتنی بر بهینه سازی کولونی مورچه ای را در زیر می آوریم (47)

1. TrainingSet = {all training cases};
2. Initialize population of each skill in colony;
3. DiscoveredRuleList = []; /*rule list is initialized with an empty list*/
4. While (TrainingSet > Max_uncovered_cases)
5. j = 1; /* convergence test index */
6. s = 1; /* Skill index */
7. Initialize all trails with the same amount of pheromone;
8. MultiClassRule = []; /*a multi class rule buffer is initialized with an empty list */
9. Evaluate maximum skills(i.e. MaxSkill) of TrainingSet;
10. While (s<MaxSkill)
11. Calculate starting ant index (St(s)) for s-th skill;
12. t = St(s);
13. Calculate ending ant index (i.e. EndIndex = MaxSkill(s) + St(s) for s-th skill;
14. Repeat
15. Antt starts with an empty rule and incrementally constructs a classification rule Rt by adding one term at a time to the current rule;
16. Prune rule Rt;
17. Update the pheromone of all trails by increasing pheromone in the trail followed by Antt (proportional to the quality of Rt) and decreasing pheromone in the other trails(simulating pheromone evaporation);
18. IF (Rt is equal to Rt-1) /*update convergence test*/
19. Then j = j+1;
20. Else j = 1;

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

21. End IF
22. $t = t+1$;
23. Until ($t \geq \text{No_of_ants}$) OR ($j \geq \text{No-rules_converg}$) OR ($t \geq \text{EndIndex}$)
24. Generalize and Identify intervals for continuous attributes;
25. Update all partial paths by means of discovered intervals;
- 29". Choose the best Rbest among all rules R_t constructed by all ants of s ;
27. add rule Rbest to MultiClassRule as a single objective rule;
28. $s = s+1$
29. END While
30. Add single objective rules from MultiClassRule to DiscoveredRuleList;
31. TrainingSet = TrainingSet – {set of cases correctly covered by MultiClassRule};
32. END While

توضیح این روش بصورت خط به خط در زیر آورده شده است:

خط 1: تمامی نمونه های آموزشی بایستی به مجموعه ی آموزشی انتقال یابد (خواندن مصادیق آموزشی)

خط 2: بایستی جمعیت هر گروه از مهارت های مورچه های متخصص مشخص گردد. بعبارت دیگر این گام معادل مشخص ساختن ساختار تخصصی یک کولونی است.

خط 3: خالی کردن و آماده سازی مجموعه ای که برای نگه داری دانش ذخیره شده استفاده می گردد تا آمادگی لازم را برای فرایند داده کاوی جدید حاصل شود. دانش استخراجی به شکل لیستی از قوانین نگهداری می گردد.

خط 4: شروع فرایند کشف دانش بصورت تکراری. بدنه اصلی کشف دانش در بین خط های 4 تا 32 قرار دارد که بصورت یک حلقه while پیاده سازی شده است. که این حلقه قوانین کلاس بندی را کشف می کند. در این الگوریتم می توان دو سطح متمایز را تشخیص داد: سطح مورچه ها (عامل ها) و سطح کولونی (سیستم چند عامله). قسمتهای داخل حلقه (Repeat –Until) با سطح مورچه ها و راه حل های جزئی آنها از مسئله در ارتباط است و بقیه ی قسمتهای داخلی حلقه ی while با سطح کولونی (سیستم چند عامله) در ارتباط است.

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

هر قانون کشف شده به لیستی از قوانین کشف شده اضافه می گردد و مصادیق آموزشی که بدرستی توسط این قانون پوشش داده شده اند (یعنی آن مواردی که مقدم قوانین را ارضاء می کند و کلاس پیش بینی شده در بخش پی آمد آن قانون را دارد) از مجموعه ی آموزشی حذف می گردند. این پردازش تا هنگامی که تعداد مصادیق آموزشی پوشش نیافته از حد آستانه ای که توسط کاربر مشخص می گردد و $Max_uncovered_cases$ نامیده می شود ، بیشتر باشد ، بصورت تکراری ادامه می یابد.

هر تکرار از حلقه Repeat-Until (سطح عوامل) متشکل از سه گام است: ساخت قانون ، هرس قانون¹ و بروز رسانی فرومون ، که به جزئیات آن در ادامه خواهیم پرداخت.

خط 5: مقدار دهی اولیه به متغیری خاص جهت آزمون همگرایی مورچه ها.

خط 6: مقدار دهی اولیه به متغیری خاص بعنوان اندیس مهارت و تخصص مورچه ها. متغیر S، شماره ی تخصص جاری را که در آن فعالیت ساخت قانون صورت می گیرد مشخص می سازد.

خط 7: مقدار دهی اولیه فرومون به تمام فضای حالت و مسیر های ممکن برای مورچه ها در همه ی تخصص ها. در این گام ابتدا جدول ساخته شده ، برای کار توسط مورچه ها مورد مقدار دهی اولیه ی فرومون قرار می گیرد. این مقدار دهی اولیه ی فرومون توسط رابطه ی 1-4 صورت می گیرد.

$$t_{ij}(t=0) = \frac{1}{\sum_{i=1}^a b_i}$$

رابطه 1-4. تابع مقدار دهی اولیه برای هر $term_{ij}$ در ابتدای شروع بکار $t=0$

¹ Rule Pruning

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در هر لایه بصورت جدا از سایر لایه ها صورت می گیرد (زیرا فرض بر مبادله نکردن اطلاعات بین مهارت‌های متفاوت است). سپس بر اساس "میزان اطلاعات" موجود در هر عبارت که در گام پیش پردازش از رابطه ی لبرای هر کدام از مقادیر عبارت ها بدست آمده و فرومون احتمال نسبی به هر عبارت محاسبه می گردد.

خط 8: یک قانون چند طبقه بعنوان بافر مقدار دهی اولیه می گردد که بتوان قوانین یک طبقه ای که مورچه ها از هر تخصص بدست می آورند را در خود نگهداری نماید.

خط 9: مقدار دهی اولیه به متغیر MaxSkill، این متغیر حداکثر مهارت‌های موجود را مشخص می سازد. مقداردهی اولیه به این متغیر بر اساس محاسبه تعداد طبقه های موجود در مجموعه ی آموزشی انجام گیرد.

خط 10: شروع حلقه While داخلی که نمایش دهنده و کنترل کننده کار بر روی مهارت‌های مختلف است. با عوض شدن مقدار S در این حلقه، لایه کاری و تخصص مورد بررسی روی آن عوض می گردد.

خط 11: محاسبه مقدار اندیس اولین مورچه در مهارت S-ام. این اندیس اولیه St(s) است.

خط 12: مقدار دهی اولیه به شمارنده t. این شمارنده هر مورچه در کولونی را صرفنظر از تخصصش مشخص می سازد. بنابراین برای دسترسی به مورچه هایی از تخصص S-ام بایستی شماره مورچه آغازین از آن تخصص را به t انتساب داد.

خط 13: اندیس نهایی مورچه های متخصص از مهارت S-ام را مشخص می سازد. این مقدار کرانی را EndIndex می نامیم که از جمع کردن تعداد کلی جمعیت مورچه ها در تخصص S-ام با اندیس آغازین مورچه ها بدست می آید.

$$\text{EndIndex} = \text{MaxSkill}(s) + \text{St}(s)$$

خط 14: ایجاد حلقه ای برای شبیه سازی کار مورچه های متخصص در یک لایه خاص (عوامل از یک مهارت) را بعهدہ دارد. ناحیه قرار گرفته در بین 14 تا 23 سطح عوامل متخصص در الگوریتم را نشان میدهد.

خط 15 و 16 و 17: ساخت قوانین (بوجود آوردن مسیر ها)، هرس قوانین و بروز رسانی فرومون در عبارت های بکار رفته بر روی تخصص S-ام از مهارت‌های موجود کولونی در حال انجام است. بعبارت دیگر تمام فعالیت های ذکر شده بر روی لایه S-ام در سیستم در حال انجام است.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 رئوسای اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

خط 18 تا 21: آزمون همگرایی برای مورچه های متخصص از مهارت S-آم را بر عهده دارد. این خطوط از الگوریتم تعداد قوانین کشف شده یکسان را در داخل تخصص S-آم (لایه جاری) مورد شناسایی قرار می دهند . حاصل این آزمون در آرگومان از الگوریتم ذخیره می گردد تا در خطوط بعدی مورد استفاده قرار گیرد.

خط 22: اندیس مورچه ی متخصص بعدی از مهارت S-آم شکل می گیرد.

خط 23: شرایط کرانی برای خاتمه ی کار مورچه های متخصص (عوامل متخصص) از مهارت S-آم بررسی می گردد. در این شرط، حداکثر مورچه های متخصص برای مهارت S-آم در کولونی که آنرا بصورت MaxSkill(s) نمایش می دهیم، بایستی مورد بررسی قرار گیرد. عبارت دیگر بایستی کنترل کرد که مورچه های متخصص از مهارت S-آم دقیقاً به تعداد مشخص شان در کولونی ایجاد و مورد استفاده قرار گیرند.

خط 24: این خط متعلق به سطح تخصصی مورچه هاست. عبارت دیگر قوانین کشف شده توسط مورچه های متخصص توسط سطح تخصصشان (لایه کاری) مورد تعمیم قرار می گیرد.

خط 25: وظیفه ی این خط بروز رسانی قوانین جزئی (مسیر های جزئی) کشف شده توسط مورچه هاست. این خط از الگوریتم نیز مربوط به سطح تخصصی مورچه هاست . عبارت دیگر قوانین کشف شده توسط مورچه های متخصص توسط سطح تخصصشان (لایه کاری) مورد تعمیم قرار می گیرد .

خط 27: این بخش الگوریتم به انتخاب بهترین قانون اکتشافی (بهترین قانون جزئی کشف شده) می پردازد. اضافه نمودن بهترین قانون کشف شده در مهارت S-آم (لایه S-آم) به قانون چند طبقه بافر¹ انجام می گیرد.

خط 28: اضافه نمودن به شمارنده مربوط به مهارتها، این اضافه نمودن لایه کاری و تخصص مورد کار را عوض می کند.

خط 29: پایان سطح تخصص .

خط 30: قانون چند طبقه اکتشافی که حاوی بهترین قانون های یک طبقه ای کشف شده در هر تخصص است ، بایستی به مجموعه ی قوانین اضافه گردد.

خط 31: مجموعه ی آموزشی بایستی بروز رسانی گردد، به نحوی که مجموعه ی نمونه هایی که توسط قانون چند طبقه ی کشف شده پوشش داده شوند، از مجموعه ی آموزشی حذف گردند . حذف نمونه های پوشش یافته

¹ Multi Class Rule

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

باعث می‌گردد که در هر چرخه از حلقه While، کولونی با فضای مسئله‌ی جدیدی روبرو گردد. در نتیجه حتی امکان دارد که تعداد طبقه‌های موجود در مجموعه‌ی آموزشی در طی این بروز رسانی تغییر کند که این تغییرات چرخه‌های بعدی الگوریتم را متأثر می‌کند.

1-7-4- ساخت قانون:

در ابتدا مورچه‌ی t -ام (Ant_t) با قانونی خالی شروع بکار می‌کند که یعنی قانونی با هیچ عبارتی در مقدمش و در هر لحظه یک عبارت به آن اضافه می‌کند، عبارت دیگر مسیری جزئی برای هر مورچه در حال شکل گرفتن است. قانون جزئی ساخته شده جاری برای هر مورچه با مسیر جزئی جاری طی شده بوسیله‌ی آن مورچه در ارتباط است. انتخاب هر عبارت برای اضافه شدن به قانون جزئی جاری با انتخاب مسیر حرکت بطوریکه مسیر جزئی ادامه یابد مشابه است. انتخاب عبارت برای اضافه شدن به قانون جزئی جاری بستگی به دو مطلب دارد: تابع هیورستیکی وابسته به مسئله‌ی (h) و مقدار فرومون (t) که متناسب به هر عبارت است. مورچه‌ی t -ام به اضافه کردن عبارت‌ها بصورت یکی یکی به قانون جزئی جاری مربوط به خودش در هر دفعه اجرا حلقه ادامه می‌دهد تا اینکه یکی از شرایط زیر برقرار شود:

- هر عبارتی که به قانون اضافه می‌گردد باعث می‌شود آن قانون موارد کمتری را پوشش دهد، هرگاه این تعداد کمتر از حد آستانه‌ای که توسط کاربر تعریف گردیده است که در الگوریتم آنرا تحت نام $Min_cases_per_Rule$ می‌شناسیم (حداقل تعدادی از نمونه‌ها که بوسیله‌ی قانون پوشیده شود) کار خاتمه می‌یابد.

همه‌ی صفات مشخصه قبلاً بوسیله‌ی آن مورچه استفاده شده باشد بنابراین صفات مشخصه‌ی بیشتری برای اضافه کردن به مقدم قانون وجود نداشته باشد. قابل توجه است که هر صفت مشخصه بیشتر از یکبار نمی‌تواند در هر قانون ظاهر گردد.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

2-7-4- هرس قانون

در گام دوم این حلقه، قانون R_t که بوسیله ی مورچه ی Ant_t ساخته شده است به منظور حذف عبارت های بی ربط، بریده¹ می گردد که بعداً مورد بحث دقیق تر قرار خواهد گرفت.

در اینجا فقط باید اشاره شود که ساخته شدن عبارت های غیر مرتبط در قانون جاری براساس متغیر های اتفاقی در روال انتخاب عبارت ها و یا بر اثر دید مقطعی تابع هیورستیک که فقط یک صفت مشخصه را در هر دفعه مورد توجه قرار می دهد و از بقیه ی صفات مشخصه صرفنظر می کند، ممکن است حاصل شده باشد.

3-7-4- بروز رسانی فرومون

سومین گام در این حلقه است و میزان فرومون را در هر مسیر مورد بروز رسانی قرار می دهد که فرومون مسیری را که بوسیله ی مورچه ی Ant_t طی شده بر اساس کیفیت انتساب داده شده به قانون R_t بایستی افزایش یابد و میزان فرومون سایر مسیر ها بایستی کاهش یابد که این کاهش شبیه سازی حالت تصعید فرومون در محیط طبیعی است.

سپس مورچه بعدی شروع به ساخت قانون خودش می کند که از مقادیر جدید فرومون بعنوان راهنما ی جستجوی مربوط به خودش استفاده می کند، این فرایند (حلقه ی Repeat Until) تا رسیدن به یکی از شروط زیر ادامه می یابد:

- تعداد قوانین ساخته شده مساوی یا بزرگتر از حد آستانه ای No_of_ants که توسط کاربر مشخص می گردد باشد.
 - مورچه ی جاری Ant_t قانونی ساخته باشد که دقیقاً مشابه قانون ساخته شده توسط (No_rules_converg-1) مورچه ی قبلی باشد در حالیکه No_rules_converg نشان دهنده ی تعداد قوانینی است که برای آزمون همگرایی مورچه ها بکار رفته است.
- وقتی که حلقه ی Repeat-Until کامل شد یعنی در سطح کولونی (سیستم چند عامله) بهترین قانون در میان قانون های ساخته شده بوسیله ی تمام مورچه ها انتخاب و به فهرست قوانین کشف شده اضافه می گردد و سیستم

¹ Prune

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

چرخه ی جدیدی از حلقه ی While را بوسیله ی مقدار دهی اولیه به تمام فرمولونها در تمام مسیرها آغاز می کند

لازم بذکر است که در تعریف استاندارد از ACO در مرجع [21] جمعیتی از مورچه ها بعنوان مجموعه ای از مورچه ها تعریف گردیده که دائماً در حال ساخت راه حل هایی در بین دو بروز رسانی فرمون هستند . بر اساس این تعریف در هر جرخه از حلقه ی While الگوریتم با جمعیت یک مورچه کار می کند زیرا بروز رسانی فرمون بعد از ساخت قانون توسط مورچه است . بنابراین می توانیم بگوییم هر حلقه ی While از الگوریتم ، یک مورچه ی منفرد دارد که تکرار های زیادی را انجام می دهد . توجه به این مطلب ضروری است که تکرارهای متفاوت از حلقه While با جمعیت های متفاوتی در ارتباط است ، زیرا هر جمعیت از مورچه ها شروع به حل مسئله ی جداگانه ای می نمایند زیرا مجموعه ی آموزشی هر کدام تفاوت می کند (به دلیل حذف موارد پوشش یافته توسط قوانین کشف شده) البته در متن ، t -آمین تکرار از آن مورچه ، مورچه ی t -آم ، Ant_t نامیده شده است که به منظور توصیف ساده تر الگوریتم است .

4-7-4- ساخت قوانین طبقه بندی

از دیدگاه کشف دانش ، هسته ی این الگوریتم در گام اول حلقه ی Repeat –Until است که در آن مورچه ی جاری بصورت تکراری در هر دفعه تکرار یک عبارت به قانون جزئی خودش اضافه می نماید . در صورتیکه هر عبارت را بصورت $term_{ij}$ فرض کنیم که معنی آن عبارتی با شرطی به فرم $A_i = V_{ij}$ است که A_i صفت مشخصه ی i -آم از دامنه ی A_i . در این حال احتمال آنکه عبارت مذکور $term_{ij}$ برای اضافه شدن به مسیر جزئی جاری انتخاب گردد از رابطه ی زیر بدست می آید:

$$P_{ij} = \frac{h_{ij} \cdot t_{ij}(t)}{\sum_{i=1}^a X_i \cdot \sum_{j=1}^{b_i} (h_{ij} \cdot t_{ij}(t))}$$

رابطه 4-2. احتمال انتخاب هر عبارت برای اضافه شدن به مسیر جزئی جاری

که در این رابطه شرح متغیرها بصورت زیر است:

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

h_{ij} مقدار تابع هیوریستیک وابسته به مسئله برای عبارت ($term_{ij}$) است. هرچه مقدار h_{ij} بیشتر باشد عبارت ($term_{ij}$) برای کلاس بندی کاربردی تر است و در نتیجه بایستی احتمالش برای انتخاب شدن بیشتر باشد. تابعی که مقدار هیوریستیک وابسته به مسئله را مشخص می سازد بر اساس تئوری اطلاعات پایه ریزی گردیده و در ادامه مورد بحث قرار می گیرد.

$t_{ij}(t)$ میزان فرومون انتساب داده شده به عبارت ($term_{ij}$) در تکرار t -ام است که با مقدار فرومون موجود جاری در مکان i و j از مسیری که مورچه ی جاری پیموده، مرتبط است. هرچه کیفیت قانون ساخته شده بوسیله ی یک مورچه بهتر باشد، به میزان فرومون اضافه شده به آن قسمت از مسیر که توسط مورچه ملاقات شده اضافه می گردد. بنابراین هر چه که زمان می گذرد بهترین قطعات از مسیر طی شده - یعنی بهترین عبارت هایی (زوجهای خصیصه - ارزش) که به قانون اضافه گردیده - مقدار فرومون بیشتر و بیشتری خواهند داشت که باعث افزایش احتمال انتخاب آنها خواهد شد.

a تعداد همه ی صفات مشخصه است.

x_i مجموعه ایست که مانند پرچم عمل می کند و به ازای هر صفت مشخصه، عضو دارد. اگر صفت مشخصه ی A_i تا بحال بوسیله ی مورچه ی جاری استفاده نشده باشد عضو متناظرش در این مجموعه یک است، و در غیر اینصورت صفر خواهد بود.

b_i تعداد مقادیر دامنه ی i -امین صفت مشخصه را نشان می دهد.

عبارت ($term_{ij}$) برای اضافه شدن به قانون جزئی جاری متناسب با احتمال بدست آمده از فرمول فوق انتخاب می گردد، البته با توجه به دو محدودیت زیر:

1. صفت مشخصه ی A_i نمی تواند قبلاً در قانون جزئی جاری انتخاب شده باشد که این مطلب برای ارضای این محدودیت است که مورچه ها بایستی "بخاطر داشته باشند" که کدامیک از عبارت ها (زوج های خصیصه - ارزش) در قانون جزئی جاری اتفاق افتاده اند.

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 ثروای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

2. یک عبارت ($term_{ij}$) نمی تواند به قانون جزئی جاری اضافه گردد اگر این اضافه شدن باعث شود که آن قانون از حدی که قبلاً مشخص شده تعدادی مصادیق کمتری را مورد پوشش قرار دهد که این حد آستانه را $Min_cases_per_rule$ می نامیم.

وقتی که مقدم قانون کامل گردید ، الگوریتم بخش پی آمد را انتخاب می کند (یعنی کلاس پیش بینی شده توسط قانون انتخابی است) بطوریکه کیفیت قانون ایجاد شده بیشینه گردد . این مطلب بوسیله ی نسبت دادن کلاس غالب¹ در میان موارد پوشش داده شده توسط قانون به بخش پی آمد قانون صورت می گیرد.

5-7-4- تابع هیوریستیک

برای هر عبارت ($term_{ij}$) که توانایی اضافه شدن به قانون جاری را داشته باشد ، الگوریتم ، مقدار تابع هیوریستیک h_{ij} را که تخمینی از کیفیت این عبارت است را با توجه به توانایی آن در بهبود دادن صحت پیش بینی² قانون محاسبه می کند . این تابع هیوریستیک بر پایه ی تئوری اطلاعات (40) محاسبه می گردد. عبارت دقیق تر مقدار h_{ij} برای ($term_{ij}$) دربرگیرنده ی میزان انتروپی (مقدار اطلاعات) انتساب داده شده به آن عبارت است . برای هر ($term_{ij}$) که به شکل $A_i = V_{ij}$ است (که در آن A_i صفت مشخصه ی i - اُم است و V_{ij} مقدار j - اُم از دامنه ی A_i میزان انتروپی از رابطه ی 3-4 محاسبه می گردد): (30)

$$H(W / A_i = V_{ij}) = - \sum_{w=1}^k (P(w / A_i = V_{ij}) \cdot \log_2 P(A_i = V_{ij}))$$

رابطه ی 3-4. میزان انتروپی عبارت $term_{ij}$

که در این رابطه :

W صفت مشخصه ی کلاس (یعنی صفت مشخصه ای که دامنه ی آن شامل طبقه هایی برای پیش بینی است .
 K تعداد طبقه هاست .

¹ Majority class

² Predictive accuracy

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

$P(w/A_i = V_{ij})$ احتمال تجربی مشاهده ی کلاس w است بشرط مشاهده ی عبارت $term_{ij}$ (یعنی همان $A_i = V_{ij}$).

نکته: برای محاسبه ی احتمال تجربی کافیت که رابطه ی $|T_{ij}| / freq_{ij}^w$ محاسبه گردد که در آن $|T_{ij}|$ تعداد همه ی مصادیق در پارتیشن T_{ij} است (پارتیشنی شامل آن نمونه ها و مواردی که صفت مشخصه A_i مقدار V_{ij} را دارد) و $freq_{ij}^w$ تعداد مصادیقی است که در پارتیشن T_{ij} کلاس w دارند. (38)

هرچه مقدار انتروپی $H(W/A_i = V_{ij})$ بیشتر باشد نشان دهنده ی آن است که طبقه ها بطور یکنواخت تری¹ توزیع شده اند و بنابراین بایستی مورچه ی جاری با احتمال کمتری آن عبارت $(term_{ij})$ را برای قانون جزئی خودش مورد انتخاب قرار دهد.

برای فرمول 1، نرمال بودن و بهنجاری میزان تابع هیورستیک ایده آل است و بمنظور پیاده سازی این نرمال سازی از این حقیقت استفاده می کنیم که مقدار تابع $H(W/A_i = V_{ij})$ در محدوده $0 \leq H(W/A_i = V_i) \leq (\log_2 k)$ تغییر می کند که در آن k ، تعداد طبقه ها را نشان می دهد. بنابراین تابع هیورستیک نرمال شده بر اساس تئوری اطلاعات به شکل فرمول 4-4 است:

$$h_{ij} = \frac{\log_2 k - H(W/A_i = V_{ij})}{\sum_{i=1}^a X_i \cdot \sum_{j=1}^{b_i} (\log_2 k - H(W/A_i = V_{ij}))}$$

رابطه 4-4. تابع هیورستیک نرمال شده برای عبارت $term_{ij}$

که a و X_i و b_i همان معانی گفته شده در رابطه ی قبل را دارند.

توجه به این نکته لازم است که مقدار انتروپی $H(W/A_i = V_{ij})$ یک عبارت $(term_{ij})$ ، صرفنظر از محتویات قانونی که عبارت در آن بکار رفته است، همواره یکسان است. بنابراین به منظور صرفه جویی در زمان محاسبه،

¹ Uniformly

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

($H(W / A_i = V_{ij})$) برای همه ی عبارات ($term_{ij}$) در گام پیش پردازش محاسبه می گردد. در رابطه ی 3 فقط دو نکته مهم بعنوان تذکر وجود دارد:

1. مقدار V_{ij} از صفت مشخصه ی A_i در مجموعه ی آموزشی حادث نشده باشد ، بنابراین $H(W / A_i = V_{ij})$ بیشینه ی خود را خواهد داشت که مقدار $\log_2 k$ خواهد بود که این مطلب معادل انتساب کمترین قدرت پیش بینی ممکن به عبارت ($term_{ij}$) است.

2. اگر همه ی موارد متعلق به یک کلاس باشند بنابراین $H(W / A_i = V_{ij})$ (مقدارش مساوی صفر خواهد شد که این مطلب مرتبط با انتساب دادن حداکثر قدرت پیش بینی به عبارت ($term_{ij}$) است.

در الگوریتم کشف دانش مبتنی بر کولونی مورچه ، اندازه انتروپی با بروز رسانی فرومون بصورت توأم با هم بکار می رود . کاربرد توأم اندازه یا انتروپی و بروز رسانی فرومون باعث مقاوم شدن این الگوریتم در برابر خطا می گردد و همچنین باعث می شود که الگوریتم ، کمتر در دامهای بهینه ی محلی در فضای جستجو گرفتار آید، زیرا بازخورد ارائه شده بوسیله ی بروز رسانی فرومون به تصحیح بعضی از خطاهای رخ داده شده توسط دید محدود اندازه ی انتروپی کمک می نماید. توجه به این نکته ضروری است که اندازه ی انتروپی ، اندازه ی هیوریستیک و محلی است که تنها به یک صفت خاصه در هر لحظه توجه می نماید و بنابراین به مسئله ی تعامل¹ بین صفت های مشخصه ، حساس است .

در مقابل ، بروز رسانی فرومون تمایل به تطبیق بهتری با تعامل بین صفات مشخصه دارد ، زیرا بروز رسانی فرومون مستقیماً بر پایه ی کارائی قانون بعنوان یک موجودیت استوار است. یعنی فرومون ، مستقیماً تعامل میان تمام صفات مشخصه ی موجود در قانون را در نظر می گیرد.

6-7-4- هرس قوانین طبقه بندی

هرس ی قطع قوانین² تکنیکی عمومی در زمینه ی کشف دانش است (29). هدف اصلی از قطع قوانین ، حذف عبارتهای بی ربطی است که بطور زائد به قانون اضافه گردیده است . قطع قوانین ، توانایی بالقوه ی افزایش قدرت

¹ Interaction

² Rule Pruning

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

پیش بینی قوانین را داراست که اینکار را از طریق ممانعت بیش از حد¹ قانون انجام می دهد . علت دیگر برای قطع قوانین آنست که اینکار سادگی قانون را افزایش می دهد، زیرا هرچه قانون کوتاهتر باشد دارای درک بالاتری توسط کاربر خواهد بود.

بعد از اتمام کار یک مورچه در ساخت قانون مربوط به خودش ،پردازش مربوط به قطع قانون فراخوانی می گردد. این راهکار برای پردازش قطع قانون مشابه راهکاری است که توسط مرجع (2) بکار گرفته شده ، اما ملاک کیفیت استفاده برای قانون خیلی باهم متفاوت هستند.

ایده ی اصلی در قطع قوانین ، حذف یک عبارت از قانون در هر دفعه است که بصورت مکرر مادامیکه این پردازش کیفیت قانون را افزایش دهد، اجرا می گردد.

7-4-7- قوانین انتقال و بروز رسانی فرومون

همانطوریکه گفته شد ، هر عبارت ($term_{ij}$) ، با قطعه ای از مسیری که توسط مورچه طی گردیده ، در ارتباط است .

در هر تکرار حلقه While از الگوریتم (در سطح کولونی) همه ی عبارات ($term_{ij}$) به مقدار اولیه ای از فرومون مقدار دهی می شوند .مورچه ها ترجیح می دهند به سوی معنایی حرکت کنند که بوسیله ی مسیر هایی با مقدار بالای فرومون متصل شده باشند . بخشی از فرومون بر روی همه ی مسیرها بخار می شود و سپس هر مورچه مقداری فرومون بر روی مسیر ها باقی می گذارد. و مسیر هایی که به سفرهای خوب تعلق دارند مقدار بیشتری فرومون دریافت می کنند و این فرایند تکرار می شود.به عبارتی این شبیه یک برنامه ی تقویت یادگیری است که در آن راه حل های بهتر ، تقویت بالاتری دریافت می کنند

فرومون موجود بر روی این مسیر ها نقش یک حافظه ی بلند مدت توزیع شده را بازی می کنند .

پس بعد از اینکه همه عبارات به مقدار اولیه ای از فرومون مقدار دهی شدند ، اولین مورچه ، جستجوی خودش را در حالی شروع می کند که همه مسیر ها مقدار دهی اولیه شده اند . دراین الگوریتم همه ی مسیر ها یک مقدار

¹ Over fitting

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

فرومون را دارا هستند، که این مقدار اولیه فرومون که در هر موقعیت مکانی از مسیر بر جای گذاشته شده با تعداد مقادیر همه ی صفات مشخصه نسبت عکس دارد که با رابطه 6 تعریف گردیده است.

هر جائیکه یک مورچه ساخت قانون خودش را تمام کرده باشد، آنگاه مقدار فرومون در همه ی قسمتهای تمام مسیرها باید بروزرسانی گردد. این بروز رسانی فرومون بوسیله دو ایده اصلی پشتیبانی میگردد که عبارتند از:

1. مقدار فرومون انتساب داده شده به هر عبارت که در قانون بوسیله مورچه یافت گردیده به نسبت کیفیت قانون افزایش یابد.

2. مقدار فرومون انتسلی به هر عبارت که در قانون وقوع نیافته باشد کاهش یابد. در طبیعت فرومون باقی مانده در مسیرها بر اساس عوامل محیطی تضعیف و تصعید می گردد که این نرخ خود مسئله مهمی در پویایی الگوریتم بهینه سازی مورچه هاست.

افزایش مقدار فرومون انتساب داده شده به هر عبارت که در قانون بدست آمده توسط یک مورچه استفاده شده باشد، مرتبط با افزایش دادن احتمال انتخاب شدن آن عبارت توسط سایر مورچه ها در آینده است که بایستی مرتبط با کیفیت آن قانون باشد. کیفیت یک قانون که با Q نمایش داده می شود بوسیله رابطه کلی زیر تعریف شده است (45)

$$Q = \text{Sensitivity} * \text{Specificity}$$

رابطه 4-5. کیفیت قانونی که با Q نمایش داده می شود

که این رابطه را به صورت زیر باز می کنیم:

$$Q = \frac{TP}{TP + FN} \cdot \frac{TN}{FP + TN}$$

رابطه 4-6. محاسبه درجه کیفیت یک قانون

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

که آرگومانهای آن بصورت زیر تعریف می گردند:

TP (True positive): تعداد مصادیقی که بوسیله قانون پوشش یافته باشند و طبقه پیش بینی شده توسط قانون را داشته باشند.

FP (False positive): تعداد مصادیقی که بوسیله قانون پوشش یافته باشند ولی طبقه پیش بینی شده توسط قانون را نداشته باشند.

FN (False negative): تعداد مصادیقی که بوسیله قانون پوشش نیافته باشند اما طبقه پیش بینی شده توسط قانون را داشته باشند.

TN (True negative): تعداد مصادیقی که بوسیله قانون پوشش نیافته باشند و طبقه پیش بینی شده توسط قانون را نداشته باشند.

مقدار Q در محدوده $0 \leq Q \leq 1$ تغییر می کند. هر چه مقدار Q بیشتر باشد کیفیت قانون بالاتر است. بروز رسانی فرومون برای هر عبارت بر اساس رابطه 4-7 برای همه ی عبارتهایی که در قانون واقع شده باشند انجام می گیرد:

$$t_{ij}(t+1) = t_{ij}(t) + t_{ij}(t).Q, \forall i, j \in R$$

رابطه 4-7 - مشخص کننده نحوه ی بروز رسانی فرومون برای هر عبارت

که در آن R مجموعه عبارت هایی است که در قانون ساخته شده بوسیله مورچه در تکرار t -ام (مورچه t -ام) از الگوریتم وجود دارند. بنابراین مقدار فرومون های همه ی عبارت های موجود در قانون کشف شده بوسیله مورچه جاری، به نسبت مقدار فرومون جاری آن عبارت ها و کیفیت اخذ شده نهایی از قانون افزایش می یابند. همانطوری که قبلاً اشاره شد میزان فرومون انتساب داده شده به هر عبارت که در قانون کشف شده ی مورچه ی جاری حضور ندارد بایستی به منظور شبیه سازی تصعید و تبخیر فرومون در کولونی ها ی مورچه های واقعی، کاهش یابد.

تأثیر تصعید فرومون برای عبارتهای استفاده نشده بوسیله ی نرمال سازی هر مقدار فرومون t_{ij} بدست می آید. این نرمال سازی بوسیله تقسیم مقدار هریک از مقادیر t_{ij} بر جمع t_{ij} به ازای همه ی i, j ها انجام می گیرد. برای

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

دیدن چگونگی تصعید فرومون ، به این نکته توجه داشته باشید که فقط میزان فرومون آن عبارت هایی که بوسیله مورچه استفاده شده توسط رابطه ی 12 افزایش می یابد. بنابراین در زمان نرمال سازی میزان فرومون یک عبارت استفاده نشده بوسیله ی تقسیم مقدار جاری آن بر کل حاصل جمع فرومون های همه ی عبارت ها ، پیاده سازی می شود که باعث افت مقدار فرومون برای آن عبارت استفاده نشده خواهد شد.

فصل پنجم: ابهام زدایی معنی واژه

1-5- ابهام زدایی واژه با استفاده از کولونی مورچه

برای ابهام زدایی از معنی واژه ، یکی از اقدامات اصلی و اولیه در ابهام زدایی با استفاده از این روش ، تهیه و آماده سازی یک پیکره مناسب می باشد. پیکره ای که در اینجا از آن استفاده شده است بخش کوچکی از پیکره خانم دکتر طیه موسوی میانگه است. این پیکره حاوی 4 میلیون و هشتصد و شصت واژه بوده اما بدلیل اینکه برای اعمال الگوریتم مجبور هستیم که پیکره موجود را به شکل جفت های خصیصه - ارزش بکار ببریم ، تنها از قسمت کوچکی از این پیکره که حاوی 100 واژه مبهم است استفاده شده است .

برای مثال جمله زیر با واژه مبهم article را در نظر می گیریم:

Some of the articles of the constitution need to be amended.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

قسمتی از برون‌داد این پیکره برای این واژه با چهار معنی مختلف که به شکل جفت‌های خصیصه - ارزش تبدیل شده است را در جدول 5 می‌بینیم:

جدول 5. قسمتی از پیکره به شکل جفت‌های خصیصه - ارزش

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
1	Article	مقاله	constitution	قانون اساسی	0.0019
2	Article	مقاله	amend	اصلاح کردن	0.00049
3	Article	مقاله	constitution	ساختمان	0.00029
4	Article	مقاله	constitution	تشکیل	0.00012
5	Article	ماده	constitution	قانون اساسی	0.54
6	Article	ماده	constitution	ساختمان	0.00028
7	Article	ماده	constitution	تشکیل	0.0052
8	Article	ماده	amend	اصلاح کردن	0.006

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
9	Article	کالا	constitution	قانون اساسی	0.0005
10	Article	کالا	constitution	ساختمان	0.00014
11	Article	کالا	constitution	تشکیل	0.00033
12	Article	کالا	amend	اصلاح کردن	E-5.89 05

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شرایع عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

13	Article	حرف تعریف	constitution	قانون اساسی	NaN
14	Article	حرف تعریف	constitution	ساختمان	NaN
15	Article	حرف تعریف	constitution	تشکیل	NaN
16	Article	حرف تعریف	amend	اصلاح کردن	NaN
17	Article	قانون اساسی	amend	اصلاح کردن	0.029
18	constitution	ساختمان	amend	اصلاح کردن	0.0008
19	constitution	تشکیل	amend	اصلاح کردن	0.0001
20	constitution	قانون اساسی	amend	ترمیم کردن	0.00019
21	constitution	ساختمان	amend	ترمیم کردن	0.01
22	constitution	تشکیل	amend	ترمیم کردن	0.000002
23	constitution	قانون اساسی	amend	بهبود یافتن	0.0022
24	constitution	ساختمان	amend	بهبود یافتن	0.00029
25	constitution	تشکیل	amend	بهبود یافتن	0.00013

حال با توجه به مطالبی که در بخش 1-4 مطرح کردیم ، می توان قوانین را بصورت زیر داشت:

مقدم : اگر واژه مبهم Article باشد و در حالت دستوری اسم باشد و واژه همبسته آن با فاصله حداکثر 8

واژه از دوطرف در معنی "قانون اساسی" بکار رفته باشد

تالی : آنگاه معنی واژه مبهم "ماده" خواهد بود.

نکته دیگری که برای ساده کردن کار در نظر گرفته شده است ، فرض بر این است که حالت دستوری واژه از قبل تشخیص داده شده است ، یعنی میدانیم که واژه مبهم در حالت دستوری اسم قرار دارد و حال با اعمال کولونی مورچه می خواهیم تشخیص دهیم که کدامیک از معانی مختلف این واژه در حالت اسم ، به عنوان معنای درست برای آن واژه تشخیص داده می شود.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

اکنون به سراغ روش کار می‌رویم.

در این روش، ابتدا در پیکره زبان فارسی کلماتی که بصورت دوتایی هستند را جستجو می‌نماییم. سپس با هم آبی آنها را در جدولی بنام جدول با هم آبی ذخیره می‌کنیم. بدین معنا که جدولی را بصورت زیر ایجاد می‌کنیم.

جدول 6. با هم آبی دو واژه

کلمه فارسی A	
b	a
d	c

که در این جدول a، تعداد جملاتی در پیکره ی زبان فارسی است که در آن جملات، کلمه ی A و کلمه B با هم در آن جمله ظاهر شده اند، b، تعداد جملاتی است که در آن جملات، کلمه A به تنهایی و بدون کلمه B دیده شده است؛ c، تعداد جملاتی است که در آن کلمه A به تنهایی و بدون کلمه B دیده شده است و بالاخره d تعداد جملاتی است که در آنها نه کلمه ی A دیده شده و نه کلمه B.

حال برای هر جدول، میزان همبستگی¹ آن دو کلمه را با استفاده از فرمول زیر محاسبه می‌کنیم:

$$AS = \frac{(ad - bc)^2}{(a + b)(a + c)(b + d)(c + d)}$$

رابطه 1-5. فرمول محاسبه امتیاز با هم آبی بین دو کلمه

¹ - Association Score

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 رئیس‌جمهوری عالی اطلاع‌رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در این صورت همواره AS عددی بین صفر و یک خواهد بود. بعداً از این عدد برای مقدار دهی اولیه فرومون روی یال‌های انتهایی استفاده می‌کنیم. (48)

اکنون در برخورد با یک جمله دارای چند واژه ی مبهم، ابتدا گرافی برای آن واژه‌های مبهم، بشکلی که کلیه معانی مختلف آن واژه‌ها در جمله را پوشش دهد ایجاد کرده (به مثال توجه نمایید) در مرحله ی بعد با توجه به تعداد معانی مختلف هر واژه برای آن جمله مورچه ی متخصص تعریف می‌کنیم و برای هر تخصص نیز یک ماکزیمم سطح فرومون تعریف می‌کنیم یعنی اگر میزان فرومون از a تا b بود پس در حیطه تخصص مورچه‌های متخصص X قرار دارد و....). و با مراجعه به پیکره ی زبان مقصد احتمال وقوع هر معنی را بدست آورده و یالهای مربوطه را مقداردهی اولیه ی فرومون می‌نماییم. در مرحله ی بعد، با استفاده از جدول وابستگی و میزان همبستگی محاسبه شده، یالهای مربوطه را نیز مقدار دهی کرده و با مراجعه به پیکره ی زبان مقصد احتمال وقوع هر معنی را بدست آورده و یالهای مربوطه را مقداردهی اولیه ی فرومون می‌نماییم.

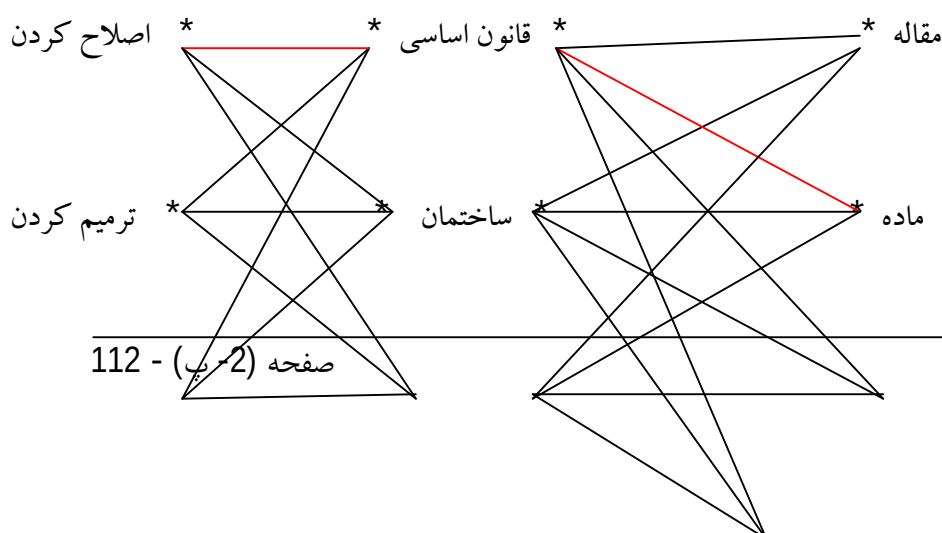
همان جمله فوق را در نظر می‌گیریم که در آن علاوه بر واژه article واژه‌های constitution و amend نیز مبهم می‌باشند:

Some of the articles of the constitution need to be amended.

article: مقاله / ماده / کالا / حرف تعریف

constitution: قانون اساسی / ساختمان / تشکیل

amend: اصلاح کردن / ترمیم کردن / بهبود یافتن



 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	

* بهبود یافتن

* تشکیل

* کالا

* حرف تعریف

شکل 9. مثالی از نحوه ایجاد گراف برای واژه مبهم

بدین ترتیب بعد از محاسبه ی مقدار AS برای هر یک از یالهای گراف فوق را از روی پیکره، و اعمال این الگوریتم می بینیم که مورچه در مسیر یالهایی که در شکل فوق با رنگ قرمز نشان داده شده اند، همگرا می شوند. گام بعدی در این روش، در نظر گرفتن رابطه ایست که به کمک آن بتوان میزان فرومون را در هر مسیر مورد بروز رسانی قرار داد. یعنی رابطه ای که قادر باشد فرومون مسیری را که بوسیله ی مورچه ی Ant_i طی شده بر اساس کیفیت انتساب داده شده به قانون R_i افزایش داده و میزان فرومون سایر مسیر ها کاهش یابد و این افزایش و کاهش شبیه سازی حالت تولید بیشتر و یا تصعید فرومون در محیط طبیعی است.

همانطور که قبلاً توضیح داده شد، هسته ی این الگوریتم حلقه تکراری است که در آن مورچه ی جاری بصورت تکراری در هر دفعه تکرار یک عبارت به قانون جزئی خودش اضافه می نماید. و دیدیم که احتمال آنکه عبارت مذکور $term_{ij}$ برای اضافه شدن به مسیر جزئی جاری انتخاب گردد از رابطه ی 2-4 بدست می آید.

تابعی که مقدار هیوریستیک وابسته به مسئله را مشخص می سازد بر اساس تئوری اطلاعات پایه ریزی گردیده است. پس از اعمال الگوریتم، به واژه ی مبهم، از بین معانی مختلف، معنی مناسبی متناظر می شود.

2-5- نتایج و بررسی کارایی الگوریتم کولونی مورچه در ابهام زدایی معنی واژه

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

الگوریتم پیاده سازی شده توسط این پروژه ، که نتایج آن در این فصل ارائه می گردد همگی در محیط ++VC ویرایش 6 بر روی Windows XP Proftional تهیه و اجرا گردیده اند. این الگوریتم بر روی سخت افزار Pentium IV پیاده سازی و تست شده اند.

در ابتدا لازم است شرایط آزمون را مشخص سازیم.شرایطی مانند نحوه ی کلاس بندی بوسیله قوانین کشف شده ، تنظیم پارامترهاو... که در زیر ارائه می گردند.

3-5- نحوه کلاس بندی بوسیله قوانین کشف شده

در الگوریتم ذکر شده به منظور کلاس بندی کردن مجموعه مصادیق¹ تست جدید¹ که در خلال آموزش دیده نشده بودند ، قوانین کشف شده به ترتیبی که کشف شده اند بکار می روند. یعنی اولین قانونی که مورد جدید را پوشش دهد برای کلاس بندی آن قانون انتخاب می گردد. امکان دارد که هیچ قانونی در لیست قوانین کشف شده برای پوشش مورد جدید وجود نداشته باشد ، در این حالت مورد جدید را تحت قانون همیشگی² کلاس بندی می کنیم .یعنی کلاس غالب را برای مجموعه موارد آموزشی کشف نشده ، مورد پیش بینی قرار می دهیم.

3-5- مجموعه آموزشی

کارایی الگوریتم کولونی مورچه بر روی مجموعه آموزشی به عنوان پیکره ای حاوی 100 واژه مبهم مورد آزمایش قرار گرفته است .مشخصه های این مجموعه آموزشی در جدول زیر خلاصه شده است.اولین ستون جدول نام مجموعه داده هاست. بقیه ستون ها به ترتیب عبارتند از: تعداد مصادیق (نمونه های موجود) و تعداد کلاس های مربوط به مجموعه داده ها.

¹ Test Set

² Default

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

جدول 7. مشخصه‌های مجموعه آموزشی

تعداد کلاس‌ها	تعداد مصادیق	مجموعه داده‌ها
6	287	Sub_Coppus article
3	263	Sub_Corpus Constitution
4	300	Sub_Corpus Conflict
4	230	Sub_Corpus Figur

5-5- تنظیم پارامترهای الگوریتم

تعداد مورچه‌ها (No_of_ants): در این الگوریتم تعداد مورچه‌ها بصورت غیر مستقیم تأثیر دارد. بعبارت دیگر وقتی تعداد مورچه‌ها را اضافه می‌نماییم در اصل به تعداد مورچه‌های زیر مجموعه یکی از (حداقل یکی و حداکثر همه) مهارت‌ها افزوده می‌گردد. بعبارت دیگر $MaxSkill(i)$ (به ازای i نامشخص) رشد خواهد داشت. با رشد کردن تعداد جمعیت متخصص در یک مهارت از مهارت‌های موجود در کولونی در اصل شانس یافتن بهترین قانون یک طبقه‌ای در آن مهارت افزوده می‌گردد. که باز هم با این افزایش سیستم به سمت دقت بالاتر و کند شدن می‌رود.

حداقل تعداد مصادیق در هر قانون (Min_cases_per_rule): هر قانون یک طبقه‌ای بایستی حداقل، این تعداد از مصادیق را پوشش دهد که باعث تقویت نظر تعمیمی قانون کشف شده گردد. این مطلب باعث می‌گردد که از عدم تطابق قانون با مجموعه داده‌ها جلوگیری شود.

تعداد بیشینه موارد کشف نشده در مجموعه آموزشی (Max_uncovered_cases): پردازش کشف قوانین بطور تکراری اجراء می‌گردد و شرط کرانی آن همین مقدار است. بعبارت دیگر این پردازش، تا هنگامی که تعداد مصادیق مجموعه آموزشی که بوسیله هیچکدام از قوانین کشف شده پوشش نیافته‌اند کوچکتر از این حد آستانه شود ادامه می‌یابد. یعنی تعداد نمونه‌های مجموعه آموزشی اگر از این مقدار کمتر گردد کل عمل ساخت قوانین متوقف می‌گردد.

 دانشگاه گیلان	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

تعداد قوانین استفاده شده برای آزمون همگرایی مورچه‌ها (No_rules_convr) در این الگوریتم اگر مورچه‌ای از مهارت i -ام قانونی یک طبقه‌ای ساخته باشد که دقیقاً مشابه با قانون یک طبقه‌ای ساخته شده توسط (No_rules_convr-1) تعداد از مورچه‌های متخصص قبلی در مهارت S -ام (لایه S -ام) باشد، سیستم نتیجه می‌گیرد که مورچه‌های مهارت S -ام به یک قانون همگرا شده‌اند (به یک مسیر در لایه S -ام). در این حالت تکرار فعلی از حلقه Repeat_Until (سطح مورچه‌های متخصص از مهارت S -ام) متوقف می‌گردد و چرخه دیگری از حلقه while داخلی را شروع می‌کند. که معادل شروع بکار بر روی لایه جدید یا مهارت جدیدی است. در نتایج بعدی مقادیر زیر بعنوان حدود آستانه برای الگوریتم انتخاب شده‌اند:

Min_cases_per_rule=10

Max_uncovered_case=10

No_rules_convr=10

No_of_ants=3000

در این الگوریتم، در همه حالات تعداد کلی مورچه‌ها بصورت معادل بین تخصص‌های مختلف تقسیم شده است. مثلاً اگر تعداد کلاس‌ها 3 تا باشد و تعداد مورچه‌ها در کولونی برابر 3000 عدد باشد آنگاه تعداد مورچه‌ها برای هر یک از تخصص‌ها برابر 1000 فرض شده است.

$$\text{MaxSkill}(0)=\dots=\text{MaxSkill}(i)=\dots=\text{MaxSkill}(k)=(\text{No_of_ants}/k)$$

رابطه 2-5- تعداد مورچه‌ها در تخصص‌های مختلف

در این رابطه k تعداد کلاس‌های موجود در مجموعه آموزشی است.

6-5- نتایج کسب شده

نتیجه کسب شده از پیاده‌سازی الگوریتم کولونی مورچه روی پیکره زبان مقصد در جدول زیر آورده شده است. این جدول نشان دهنده صحت پیش‌بینی قوانین کشف شده حاصل از الگوریتم کولونی مورچه به کمک کشف دانش چند طبقه روی پیکره زبان مقصد است. اعداد بعد از علامت $(\pm)$ انحراف از معیار¹ می‌باشد که بعد از هر کدام از صحت‌های پیش‌بینی دیده می‌شود.

¹ Standard Deviation

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیر پروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیر پروژه: پیک‌متن فارسی — 2 — پ	

جدول 8. صحت پیش بینی قوانین کشف شده حاصل از الگوریتم کولونی مورچه
به کمک کشف دانش چند طبقه روی پیکره زبان مقصد

صحت پیش بینی به درصد	مجموعه داده ها
83%±3	Sub_Coppus article
86±2/9	Sub_Corpus Constitution
87±۲	Sub_Corpus Conflict
۸۶	Sub_Corpus Figur

4-1-5- بررسی نتایج کسب شده

این بخش به ارزیابی پیچیدگی محاسباتی¹ الگوریتم کشف دانش چند طبقه مبتنی بر کولونی مورچه ها می پردازد. این تحلیل به سه بخش تقسیم گردیده است:

- 1- پیش پردازش
- 2- یک چرخه از حلقه While بیرونی
- 3- کل چرخه مربوط به حلقه While

پیچیدگی محاسباتی پیش پردازش: پیچیدگی زمانی این گام $O(n*a)$ است که در آن n تعداد مصادیق و a تعداد صفات مشخصه است.

پیچیدگی محاسباتی یک چرخه حلقه بیرونی While: هر چرخه بیرونی While از الگوریتم کولونی مورچه با مقدار دهی اولیه فرومون شروع می شود. این فعالیت بر روی هر لایه از جدول (شکل 3) انجام می دهد و بنابراین دارای پیچیدگی از مرتبه $O(a*v*s)$ می باشد که در آن a تعداد صفات مشخصه ، v تعداد مقادیر به ازاء یک صفت مشخصه و s تعداد مهارت‌های موجود (طبقه ها) در کولونی است.

¹ Computational Complexity

 دانشگاه گیلان	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

حال اگر به حلقه درونی While توجه کنیم، هر بار اجراء حلقه درونی While با اجراء چرخه Repeat-Until و بخش تعمیم و بروز رسانی معادل است پس ابتدا به بررسی چرخه Repeat-Until می پردازیم. یک تکرار از چرخه Repeat-until شامل فعالیت های یک مورچه متخصص خواهد بود. مهمترین گام انجام گرفته برای هر مورچه متخصص در این حلقه عبارتست از:

1. ساخت قانون
2. ارزیابی قانون
3. هرس قانون
4. بروز رسانی فرومون

بنابراین هر یک تکرار از چرخه Repeat-Until دارای پیچیدگی زمانی از مرتبه $O(k^3 * a * v + n)$ است که در آن k تعداد شرط های انتخابی توسط یک مورچه متخصص، a ، تعداد صفات مشخصه، و n تعداد نمونه های آموزشی (مصادیق) است.

به منظور بدست آوردن پیچیدگی محاسباتی مربوط به اجراء کامل حلقه Repeat-Until نتیجه قبلی بایستی در Z_s ضرب گردد، که تعداد مورچه های تخصص S -ام می باشد. بنابراین اجراء کامل حلقه Repeat-Until مرتبه زمانی آن $O(k^3 * [k * a * v + n * Z_s])$ می باشد.

بخش تعمیم و بروز رسانی فرومون دارای پیچیدگی زمانی از مرتبه $O(c^*)$ است که Z_s برابر با تعداد مورچه های تخصص S -ام و C تعداد خوشه های ایجاد شده برای آن صفت خواهد بود.

پیچیدگی زمانی یکبار اجرای حلقه درونی While برابر خواهد بود با حاصل جمع مطالب فوق یعنی:

$$O(Z_s * [k * a * v + n * k^3] + v^2 + s * Z_s)$$

که می توان آنرا بصورت زیر خلاصه نمود:

$$O(Z_s * [k * a * v + n * k^3 + s] + v^2)$$

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

پیچیدگی زمانی اجرای کلی حلقه درونی While برابر خواهد بود با حاصل ضرب مرتبه فوق در تعداد کلی مهارتها (طبقه ها) در کولونی یعنی:

$$O(s * z_s * [k * a * v + n * k^3 + s] + v^2)$$

چون در این الگوریتم تعداد مورچه از هر مهارت با تعداد مورچه های سایر مهارتها مساوی در نظر گرفته شده است بنابراین حاصل ضرب تعداد مهارتها در تعداد مورچه ها ی هر مهارت ، معادل تعداد کل مورچه های کولونی است. یعنی:

$$Z = s * z_s$$

بنابراین می توان پیچیدگی زمان اجرای کلی حلقه بیرونی While را بصورت زیر ساده نمود:

$$O(z * [k * a * v + n * k^3 + s] + v^2)$$

پیچیدگی زمانی یکبار اجرای حلقه بیرونی While برابر خواهد بود با مجموع پیچیدگی های مقدار اولیه و حلقه درونی While :

$$O(z * [k * a * v + n * k^3 + s] + v^2 + a + v * s)$$

پیچیدگی محاسباتی کل چرخه While: به منظور بدست آوردن پیچیدگی محاسباتی مربوط به کل

چرخه بیرونی While بایستی نتیجه قبلی را در Γ ضرب کنیم که Γ تعداد قوانین چند طبقه کشف شده است. این مقدار بسیار متغیر است و بر اساس مجموعه داده ها بسیار متغیر می باشد. بنابراین داریم :

$$O(\Gamma * \{z * [k * a * v + n * k^3 + s] + v^2 + a + v * s\} + a * n)$$

بایستی این مطلب را مورد توجه قرار داد که این پیچیدگی محاسباتی به میزان زیادی به مقدار k (تعداد شروط هر قانون) Γ (تعداد قوانین) و v (تعداد مقادیر صفات پیوسته) بستگی دارد. مقدار k و v از مجموعه داده ای به مجموعه داده ای دیگر بسیار متغیر است و به محتوای آن مجموعه داده بستگی دارد.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:		 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

2-1-6-5- محاسن بکارگیری این روش در مقایسه با روشهای موجود

در این روش با توجه به اینکه حالت دستوری واژه مشخص می شود در ابهام زدایی از معنای واژه مانند روشهای معمول استفاده از پیکره که کل پیکره مورد جستجو قرار گیرد ، در این روش بعد از اینکه حالت دستوری واژه محرز شد و با توجه به تعریف مورچه های متخصص در این پروژه ، با توجه به حالت دستوری واژه تنها قسمتی از گراف مورد پیمایش قرار می گیرد ، از طرفی با توجه به مقدار دهی اولیه فرومون مسیر های موجود روی گراف که از روی پیکره زبان مقصد استخراج می شود ، مورچه ها در پیدا کردن معنی واژه مبهم خیلی سریع تر می توانند به سمت معنی صحیح تر همگرا شوند.

نکته ی دیگر اینکه این روش مستقل از نوع زبان مقصد می باشد و به راحتی بر روی سایر زبانها نیز بعنوان زبان مقصد می تواند مورد استفاده قرار بگیرد.

7- نتیجه گیری

دراین پروژه با استفاده از پیکره زبان مقصد و بکارگیری الگوریتم کلونی مورچه ، به ابهام زدایی واژه های چند معنایی¹ زبان مبداء ، پرداخته شد.

¹ Multi-meaning word

	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

این رساله به بررسی یکی از مشکلات شاخص در رویه خودکار سازی تبدیل یک زبان نوشتاری به زبان دیگر، و ارائه راه حلی برای رفع آن می پردازد .

اصطلاح ترجمه ماشینی نامی استاندارد و رایج برای سیستم های کامپیوتری شده یا محاسباتی است که مسئول تولید ترجمه هایی از یک زبان طبیعی به زبان طبیعی دیگر با یا بدون کمک انسان باشند. در ماشین های ترجمه سعی میشود که ترجمه مناسبی برای رشته کلمات ورودی ارائه شود به گونه ای که ترجمه حاصل نه تنها معنای متن مبدأ را انتقال داده بلکه دارای سلاست کافی در زبان مقصد باشد.

ابهام واژگانی اشاره به حالتی دارد که یا یک واحد واژگانی متعلق به طبقات دستوری یا اجزاء کلام مختلف با معانی مختلف باشد و یا یک واحد واژگانی بیش از یک معنی داشته و معانی مختلف آن متعلق به یک طبقه دستوری باشند. هدف از ابهام زدایی معنی واژه ، رفع ابهام از حالت دوم است یعنی موردی که در آن معانی مختلف واژه متعلق به طبقه دستوری یکسانی باشند. ابهام زدایی معنی واژه اصطلاحی است که به استخراج معنی صحیح و مقتضی از واژه هایی که چند معنایی هستند اطلاق می شود و این کار غیر ممکن نیست چرا که اگرچه واژه ها ممکن است معانی زیادی داشته باشند، طبیعتاً تنها یکی از آن معانی مناسب بافتی است که واژه مورد نظر در آن ظاهر می شود. هدف ابهام زدایی معنی واژه ، تعیین یکی از این معناها بر اساس بافت زبانی که این واژه در آن ظاهر شده می باشد. در بسیاری از سیستم هایی که کار ابهام زدایی معنی واژه را به عهده دارند سیستم برچسب گذاری اجزاء کلام وجود دارد که مقوله دستوری واژه ها را مشخص می نماید. مثلاً در مورد یک واژه خاص در صورتیکه سیستم برچسب گذاری اجزاء کلام تشخیص دهد که این واژه اسم است آنگاه سیستم ابهام زدایی معنی واژه ، کار خود را شروع کرده و از میان معنای موجود، معانی اسمی این واژه را جدا کرده و سپس یکی را انتخاب می نماید.

ابهام زدایی معنی واژه یک امر ضروری در ترجمه ماشینی بوده و کاربرد زیادی در آن دارد چرا که سیستم ترجمه ماشینی بایستی قادر باشد تا مثلاً معنی واژه خاصی را در یک متن انگلیسی تشخیص داده و تعیین نماید که آن را به کدام کلمه معادل در فارسی ترجمه نماید.

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

در ماشین های ترجمه با بکارگیری شیوه های آماری¹، تحلیل دستوری² در سطح وابستگی ارتباطی و تحلیل معنایی³ در سطح تشخیص جملات و اطلاعات لغوی، همراه با مشخصه های دستوری در زبان مبدا، و گرفتن بازخورد از کاربران ماشین میتوان با این ابهام مقابله کرد.

روشهای مختلفی برای ابهام زدایی از معنای واژه را مورد بررسی و مطالعه قرار دادیم و با توجه به نتایج بدست آمده ما معتقدیم که استفاده از پیکره زبان مقصد در امر ابهام زدایی واژه های چند معنایی زبان مبدأ کیفیت برونداد یک سیستم ترجمه ماشینی را تا حد قابل توجهی افزایش می دهد.

لذا بطور اجمالی کولونی مورچه و الگوریتم های بهینه سازی کولونی مطالعه شد و سپس با بکارگیری الگوریتم مذکور و پیکره زبان مقصد، روشی برای ابهام زدایی از معنی واژه ارائه شد که با توجه به روشهای قبلی که بدین منظور استفاده شده است، از دقت و صحت بالاتری برخوردار است. در مدل بیشترین احتمال همیشه گزینه ای با بیشترین احتمال به عنوان گزینه درست انتخاب می شود که دقت آن نسبت به روش موسوی و دلاور خلفی نیز پایین تر است. در روش موسوی و دلاورخلفی که از اعداد تصادفی⁴ و آزمون نکویی برازش⁵ استفاده نمودند، آنها با استفاده از اعداد تصادفی برای انتخاب بهترین معادل زبان مقصد برای یک واژه مبهم از زبان مبدأ در سیستم ترجمه ماشینی استفاده کرده و روش خود را در مورد پیکره زبانی نسبتاً بزرگی در حوضه روانشناسی در زبانهای انگلیسی و فارسی به ترتیب به عنوان زبانهای مبدأ و مقصد مورد آزمایش قرار دادند. الگوریتم آنها روش توانست مشکل ابهام 604 واژه از 764 واژه مبهم یک پیکره زبانی انگلیسی را با استفاده از متن معادل آن در زبان فارسی حل نماید که بدین ترتیب صحت این مدل به 79% رسید (12). روش موسوی و دلاور خلفی نسبت به روش انتخاب بر اساس محتمل ترین حالت صورت می گیرد دقیق تر و علمی تر است

اما در این روش با بکارگیری الگوریتم کولونی مورچه و استفاده از رابطه آماری AS بر روی یالی که مربوط با هم آیی واژه ها است نسب به روشهای قبلی از دقت بالاتری می تواند برخوردار باشد و در نتیجه در کل باعث

¹ approaches Statistical-Based

² grammatical analysis

³ semantic analysis

⁴ - random numbers

⁵ - goodness of fit test

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن-فارس — 2 — پ	

ابهام زدایی از معنای واژه با دقت بالاتری می‌شود. در نهایت دیده شد که این روش با دقت 86% درصد کارایی قابل قبولی را از خود نشان می‌دهد.

مسیر کارهایی را که می‌توان بعنوان کار تحقیقاتی آینده در زمینه این پروژه معرفی نمود عبارتند از:

- 1- راهکارهایی برای مقدار دهی اولیه فرومون
- 2- راهکارهایی برای بروز رسانی بهینه تر فرومون
- 3- بررسی راهکارهایی برای تنظیم پویای تعداد مورچه های متخصص در هر مهارت. عبارت دیگر تغییر تعداد عوامل متخصص در هر مهارت بر حسب نیاز کولونی

منابع و مراجع :

1. Frawley, W., Piatetsky-Shapiro, G., Matheus, C. J., “Knowledge Discovery in Databases: An Overview”, in Knowledge Discovery in Databases, G. Piatetsky-Shapiro and W. Frawley, (Eds.), MIT Press, 1991

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

2. M. Dorigo and L. M. Gambardella, “ Ant Colony System: A Cooperative Learning Approach to the Traveling Salesman Problem”, IEEE Trans on Evolutionary Computation1(1), 53-66, 1997.
3. M.dorigo and G.Di Caro,”The ant colony optimization meta-heuristic”, In New Ideas in Optimization,D.Corne,M.Dorigo and F.Glover Eds.London,UK:McGraw Hill,pp.11-32,1999.
4. M.Dorigo,G.Di Caro and L. M. Gambardella, “Ant algorithms for discrete optimization”,Artificial Life, vol.5,no.2,pp.137-172,1999.
5. Hutchins, W. J. and Somers, H, “An introduction to machine translation”, London, Academic Press,1993.
6. Hutchins, W. J. , “Machine translation: half a century of research and use”. Prepared for UNED Summer School at Avila, Spain. Web site: <http://ourworld.compuserve.com/homepages/WJHutchins,2003>
7. Frederking, R. and Nirenburg, S. “Tree heads are better than one”. In: Proceeding of ANLP-94, Stuttgart, Germany. (1994)
8. Freivalds, J, “ The technology of translation”. Managemant Review Vol.8. No.7, pp.48-52, (1999).
9. Heyn, M. “Integrating machine translation into translation memory systems”. In: European Association for Machine Translation Workshop proceeding, ISSCO, Geneva. pp. 111-123, (1996).
10. Hutchins, J. “Translation technology and the translator”. Proceeding of 11th Conference of the Institute of Translation and Interpreting. pp. 113-120,(1997).
11. Somers, H. “ Computers and translation: a translator’s guide”. Amsterdam, John Benjamins, (2003).
12. Mosavi Miangah, T. and Delavar Khalafi “A. Word sense disambiguation using target language corpus in a machine translation system. Literary and Linguistic Computing”, 20(2), pp. 237-249, (2005).

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 ژوئیه‌ای‌ها علی‌الحفاظ رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس 2 – پ	

13. Mosavi Miangah, T. “Morphological analysis in the system of English-Persian machine translation”. Ph. D. Dissertation, Russian Academy of Sciences, Moscow. (2002).
14. Mosavi Miangah “Acom1075 Machine translation. Multilingual Computing”, Semester2 http://www.leeds.ac.uk/acom/teaching/mlc/handouts/lecture4_handout.pdf, (2003-04).
15. Seasily, J. “ Machine translation: a survey of approaches”, Technical Report, University of Michigan, <http://www.personal.umich.edu/jseasily/home.html>, (2003).
16. Robert E. Frederking ” Corpus-Based MT and Multi-Engine MT within Nespole”, Project IST-1999-11562 , WP5 - HLT development , 30th July 2002
17. Brown, P. F.; Stephen, A.; Della pietra, V.; Della Pietra, J. and Robert L. Marcer “The mathematics of statistical machine translation: parameter estimation. Computational Linguistics”, 19(2), pp. 263-311, (1993).
18. Somers, H, Review Article: “ Example-based machine translation. Machine Translation”, 14(2), pp. 113-157. (1999).
19. Turcato, D. and Popowich, F, What is example-based machine translation. In: “Proceedings of the workshop on Example-based Machine Translation, hosted by MT-Summit VIII”, Santiago de Compostela, Spain, September, 18-22, (2001).
20. J. Casillas, O.Cordon, F. Herrera, “ Learning Fuzzy Rules Using Ant Colony Optimization Algorithms”, Supported by CICYT, PB98-1319, 1998.
21. Rayner, M. and Bouillon, P, “ Hybrid transfer in an English-French spoken language translator”, Proceeding of IA’95, Montpellier, France. <http://www.cam.sri.com/tr>, (1995).
22. Iomdin, L. and Streiter, O, “Learning from parallel corpora: Experiments in machine translation”. In: Proceeding of the Dialigue’99 International Seminar in Computational Linguistics and Application, 2, pp. 79-88, (1999).
23. Carl, M., Iomdin, L., Pease, C. and Streiter, O. “Toward a dynamic linkage of example-based and rule-based machine translation. In: Hybrid Approaches to Machine

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه:			 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه:			
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارسی — 2 — پ	

Translation”, Streiter, O. Carl, M. and Haller, J. (Eds). IAI Web-based working paper. <http://www.iai.uni-sb.de>, (2000).

24. Carl, M. and Hansen, S, “Linking translation memories with example-based machine translation”, In: Hybrid Approaches to Machine Translation. Streiter, O. Carl, M. and Haller, J. (Eds). IAI Web-based working paper. <http://www.iai.uni-sb.de>, 2000

25. <http://www.trajome.org>

26. Justeson, J. S. and Katz, S. M., “ Principled disambiguation: discriminating adjective senses with modified nouns”, Computational Linguistics, 21(1), pp. 1-28, (1995).

27. <http://www.persian-language.org/>

28. Yarowsky, D , “Word sense disambiguation using statistical models of Roget’s categories trained on large corpora”, Proceeding of 15th International Conference on Computational Linguistics, pp. 454-460, (1992).

29. M. Dorigo, A. Coloni and V. Maniezzo, “The Ant System: optimization by a colony of cooperating agents”, IEEE Transactions on Systems, Man, and Cybernetics part B, vol. 26, no. 1, pp. 29-41, 1996.

30. Bruce, R. and Wiebe, J, “Word sense disambiguation using decomposable models. Proceedings of the 32nd Annual Meeting of the Association for Computational Linguistics, (1994).

31. Ng, H. T. and Lee, H. B, “Integrating multiple knowledge sources to disambiguate word sense: an example-based approach”, Proceedings of the 34th Annual Meeting of the Association for Computational Linguistics, pp. 40-47, (1996).

32. Dagan, I. And Itai, “A. Word sense disambiguation using a second language monolingual corpus. Computational Linguistics, 20(4), pp. 563-596, (1994).

33. Hyun Ah Lee, Jong C. Park and Gil Chang Kim, “ Lexical Selection with a Target Language Monolingual Corpus and an MRD”, Korean Advanced Institute of Science and Technology, Dept. of Computer Science

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه:			 شورای عالی اطلاع رسانی
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			
	عنوان زیرپروژه:			
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

34. E. Suzuki, M.Gotoh and Y. Choki, “Bloomy Decision Tree for Multi-Objective Classification ” ,©Springer-Verlag Berlin Heidelberg,2001
35. U. M. Fayyad, G. Piatetsky-Shapiro and P. Smyth, “from data mining to knowledge discovery: an overview”,In: Advances in Knowledge Discovery & Data Mining, U.m. Fayyad, G. Piatetsky-Shapiro, P. Smyth and R. Uthurusamy(Eds) Cambridge, MA:AAA/MIT, pp.1-34,1996.
36. P.L. Carbone, “Data Mining or Knowledge Discovery in Databases”, An Overview© MITRE Corporation, 1997.
37. G. Di Caro, “An Introduction to swarm Intelligence Issues”,IDSIA, USI/SUPSI,Lugano(CH),2002.
38. M.Wooldridge and N. R. Jennings, “ Intelligent agents: Theory and practice”, The Knowledge Engineering Review, 10(2): 115-152, 1995.
39. U. Fayyad, G. P. Shapiro, and P. Smyth, “From Data Mining to Knowledge Discovery in Databases”, AI Magazine, 37-54, Fall 1996.
40. E.Bonabeau, M. Dorigo and G. Theraulaz, Swarm Intelligence: “From Natural to Artificial Systems”. New York, NY: Oxford University Press, 1999.
41. R.S. Parpinelli, H.S.Lopes and A.A.Freitas, “Data Mining with an Ant Colony Optimzation Algorithm”,IEEE Trans on Evolutionary Computation, special issue on Ant Colony Algorithms,6(4), 2002.
42. R.S. Parpinelli, H.S.Lopes and A.A.Freitas, “An Ant Colony Algorithm for Classification Rule Discovery”, Data Mining: a Heuristic Approach, 191-208, Idea Group Publishing ,London ,2002
43. D. A. Holway, A. V. Suarez, “Colony-structure variation and inter-specific competitive ability in the invasive Argentine ant”,Oecologia(2004),Received: 28 May 2003 / Accepted: 17 September 2003 / Published online: 18 October 2003 @Springer Verlag 2003.
44. Hedden, T. D. “Machine translation: a brief introduction”. <http://www.hedden.org/intro-mt.html>, (1992-2003).

 دانشگاه گیلان	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
	بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی		
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی — 2 — پ	

45. Zh.I. Reznikova, "Government and nepotism in social insects: New dimension provided by an experimental approach", c Euroasian Entomological, Journal, Volume 2, Number 1, 2003.

46. دکتر علی محمدی، ابوذر جمال نیا، "الگوریتم‌های لانه مورچه و کاربرد آن در نگهداری پیشگیرانه"، اولین همایش ملی مدیریت صنعتی، 1386

47. رضا شمسایی، "الگوریتم کشف دانش چند کلاسه با استفاده از الگوریتم بهینه سازی کولونی مورچه"، پایان نامه کارشناسی ارشد، دانشگاه علم و صنعت، صفحات 60-68، 1383.

48. بهناز شب‌رو، ماشین‌های ترجمه مبتنی بر پیکره، سمینار کارشناسی ارشد، دانشگاه علم و صنعت، صفحات 10-15، پاییز 1386

49. طیبه موسوی میانگاه، "درآمدی بر ترجمه ماشینی"، انتشارات یلدا قلم، تهران، فصل سوم، صفحات 93-1386، 85

پیوست

قسمتی از برون‌داد پیکره برای چند واژه مبهم به شکل جفت‌های خصیصه – ارزش

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
1	power	اختیار	play	بازی	.0012
2	power	اختیار	play	نمایش	7.14E-05
3	power	اختیار	play	تفریح	6.7E-06

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

0001.	بازی	play	قدرت	power	4
0026.	نمایش	play	قدرت	power	5
00018.	تفریح	play	قدرت	power	6
0014.	بازی	play	حق	power	7
0006.	نمایش	play	حق	power	8
00012.	تفریح	play	حق	power	9
NaN	مشاجره	conflict	پرگو	wordy	10
NaN	مشاجره	conflict	پرلغت	wordy	11
NaN	مشاجره	conflict	لغت دار	wordy	12
0.0006	مشاجره	conflict	لفظی	wordy	13
6.77E-06	تضاد	conflict	پرگو	wordy	14
1.22E-06	تضاد	conflict	پرلغت	wordy	15
1.44E07	تضاد	conflict	لغت دار	wordy	16
8.88E-05	تضاد	conflict	لفظی	wordy	17
3.7E-7	درگیری	conflict	پرگو	wordy	18
NaN	درگیری	conflict	پرلغت	wordy	19
1.93E-07	درگیری	conflict	لغت دار	wordy	20
0.0001	درگیری	conflict	لفظی	wordy	21
NaN	ستیزه	conflict	پرگو	wordy	22

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
------	-----------	----------------	-------------	------------------	-------------------------------------

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

NaN	ستیزه	conflict	پر لغت	wordy	23
1.43E-07	ستیزه	conflict	لغت دار	wordy	24
1.2E-05	ستیزه	conflict	لفظی	wordy	25
5.77E-05	مشاجرہ	conflict	مسلح	armed	26
3.70E-6	مشاجرہ	conflict	مسلحانہ	armed	27
.00010	مشاجرہ	conflict	مجهز	armed	28
2.05E-5	تضاد	conflict	مسلح	armed	29
2.41E-5	تضاد	conflict	مسلحانہ	armed	30
0.004	تضاد	conflict	مجهز	armed	31
.016	در گیری	conflict	مسلح	armed	32
4.95E-6	در گیری	conflict	مسلحانہ	armed	33
0.00012	در گیری	conflict	مجهز	armed	34
0.00045	ستیزه	conflict	مسلح	armed	35
2.22E-5	ستیزه	conflict	مسلحانہ	armed	36
5.77E-05	ستیزه	conflict	مجهز	armed	37
0.00037	نماد	figure	سیاسی	political	38
5.33E-07	عدد	figure	سیاسی	political	39
0.0044	شخصیت	figure	سیاسی	political	40
2.57E-05	نشانه	figure	سیاسی	political	41
0.00098	آرامش	peace	نماد	figure	42
NaN	آرامش	peace	عدد	figure	43
0.00017	آرامش	peace	شخصیت	figure	44

 سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی — 2 — پ	

9.98E-05	آرامش	peace	نشانه	figure	45
.0011	صلح	peace	نماد	figure	46

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
47	figure	عدد	peace	صلح	0.00049
48	figure	شخصیت	peace	صلح	0.00048
49	figure	نشانه	peace	صلح	0.00019
50	figure	نماد	peace	دوستی	3.95E06
51	figure	عدد	peace	دوستی	7.44E-06
52	figure	شخصیت	peace	دوستی	2.87-6
53	figure	نشانه	peace	دوستی	3.09E-05
54	civil	داخلی	code	رمز	3.71E-05
55	civil	کشوری	code	رمز	2.44E-05
56	civil	غیر نظامی	code	رمز	0.00014
57	civil	مدنی	code	رمز	NaN
58	civil	داخلی	code	قانون	0.0019
59	civil	کشوری	code	قانون	3.15E-05
60	civil	غیر نظامی	code	قانون	0.00012
61	civil	مدنی	code	قانون	.00076
62	civil	داخلی	code	علامت	3.14E-05
63	civil	کشوری	code	علامت	0.0002

 وزارت فرهنگ و آموزش عالی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک-متن فارسی — 2 — پ	

9.76E-06	علامت	code	غیرنظامی	civil	64
6.54E-05	علامت	code	مدنی	civil	65
.094	جنگ	war	داخلی	civil	66
0.0057	جنگ	war	کشوری	civil	67
0.0014	جنگ	war	غیرنظامی	civil	68
0.0044	جنگ	war	مدنی	civil	69
2.84E-05	جلسه	session	بسته	closed	70
6.44E-05	جلسه	session	سری	closed	71

ردیف	واژه مبهم	معنی واژه مبهم	واژه همبسته	معنی واژه همبسته	میزان همبستگی دو واژه در پیکره (AS)
72	closed	خاتمه یافته	session	جلسه	1.67E-06
73	closed	غیرعلنی	session	جلسه	0.0076
74	closed	بسته	session	قسمت	1.15E-05
75	closed	سری	session	قسمت	2.80E-5
76	closed	خاتمه یافته	session	قسمت	7.48E-07
77	closed	غیرعلنی	session	قسمت	NaN
78	closed	بسته	session	دوره	NaN
79	closed	سری	session	دوره	4.55E-05
80	closed	خاتمه یافته	session	دوره	0.00016
81	closed	غیرعلنی	session	دوره	NaN
82	academic	علمی	exchange	بورس	9.65E-05

 وزارت آموزش عالی و تحقیقات علمی	عنوان پروژه: فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی			 شورای عالی اطلاع رسانی
	عنوان زیرپروژه: بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
	تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن‌فارس — 2 — پ	

9.53E-05	بورس	exchange	فرهنگی	academic	83
2.72E-05	بورس	exchange	دانشگاهی	academic	84
3.31E-05	ارز	exchange	علمی	academic	85
NaN	ارز	exchange	فرهنگی	academic	86
5.7E-07	ارز	exchange	دانشگاهی	academic	87
3.42E-05	مبادله	exchange	علمی	academic	88
0.0016	مبادله	exchange	فرهنگی	academic	89
0.00087	مبادله	exchange	دانشگاهی	academic	90
2.05E-07	سرعت	rate	بورس	exchange	91
.00011	سرعت	rate	ارز	exchange	92
.00011	سرعت	rate	مبادله	exchange	93
.00016	میزان	rate	بورس	exchange	94
.00099	میزان	rate	ارز	exchange	95
میزان همبستگی دوواژه درپیکره (AS)	معنی واژه همبسته	واژه همبسته	معنی واژه مبهم	واژه مبهم	ردیف
5.79E05	میزان	rate	مبادله	exchange	۹۶
.00071	نرخ	rate	بورس	exchange	97
.38	نرخ	rate	ارز	exchange	98
.071	نرخ	rate	مبادله	exchange	99
.00028	درجه	rate	بورس	exchange	100
9.67E-05	درجه	rate	ارز	exchange	101
.00042	درجه	rate	مبادله	exchange	102

 وزارت فرهنگ و آموزش عالی	 شورای عالی اطلاع رسانی		
	عنوان پروژه:		
	فاز اول طرح جامع پیکره زبان فارسی با موضوع فاز اول مطالعاتی ایجاد پیکره متنی زبان فارسی		
	عنوان زیرپروژه:		
بررسی ابعاد و لایه‌های ابهام در واژگان مشابه زبان فارسی			
تاریخ: 1388/04/7	ویرایش: 1/0	کد زیرپروژه: پیک‌متن فارس — 2 — پ	