



پژوهشگاه علوم و فناوری اطلاعات ایران
(ایرانداک)

نقش ایرانداک در گسترش خط و زبان فارسی

۲	خلاصه برای مدیران
۳	مقدمه
۶	سامانه‌ها
۱۵	طرح‌ها
۴۰	کتاب
۴۵	سخنرانی‌ها و میزگردهای علمی
۵۸	کارگاه‌ها و دوره‌های آموزشی

بنام خدا

خلاصه برای مدیران

زبان فارسی، زبان دوم عالم اسلام و کلید بخش عظیمی از ذخایر ارزشمند علمی و ادبی تمدن اسلامی، و خود از ارکان هویت فرهنگی ملت ایران است. بنابر اصل پانزدهم قانون اساسی جمهوری اسلامی ایران، زبان و خط رسمی و مشترک مردم ایران فارسی است. اسناد و مکاتبات و متون رسمی و کتب درسی باید با این زبان و خط باشد ولی استفاده از زبان‌های محلی و قومی در مطبوعات و رسانه‌های گروهی و تدریس ادبیات آنها در مدارس، در کنار زبان فارسی آزاد است. پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) از پنج دهه پیش کار مدیریت اطلاعات علم و فناوری را در ایران آغاز کرده است. خط و زبان فارسی و کاربری آن در محیط رایانه‌ای یکی از حوزه‌های مورد توجه در اساسنامه و برنامه استراتژیک پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) است. همچنین گسترش زبان فارسی برای علم در فضای مجازی یکی از مأموریت‌های ملی ایرانداک است. اسناد بالا دستی همچون نقشه جامع علمی کشور، سند چشم انداز، برنامه ششم توسعه اقتصادی، اجتماعی و فرهنگی جمهوری اسلامی و سیاست‌های کلی علم و فناوری (نظام آموزش عالی، تحقیقات، و فناوری)، همواره این حوزه را مورد توجه قرار داده‌اند. در این گزارش، چکیده‌ای از نقش ایرانداک در گسترش خط و زبان فارسی آمده است. ایرانداک برای گسترش خط و زبان فارسی تعداد ۵۲ طرح پژوهشی را انجام داده یا در دست انجام دارد که نشان دهنده جایگاه مهم این سازمان برای گسترش خط و زبان فارسی است. همچنین ایرانداک در بخش آموزش نیز تلاش‌هایی در این حوزه موضوعی انجام داده است از دیگر اقدامات ایرانداک، برگزاری ۲۰ نشست و سخنرانی در ارتباط با خط و زبان فارسی است. در این میزگردها با دعوت از کارشناسان و استادان به گفت و گوهایی انتقادی در زمینه‌های گوناگون در این حوزه پرداخته‌اند. همچنین ایرانداک ۵ عنوان کتاب در این زمینه منتشر کرده است.

شبکه‌ای و مفهومی میان واژه‌های یک یا چند حوزه از دانش بشری هستند. این نظام مفهومی با چند گونه پیوند میان واژه‌ها؛ در سازمان‌دهی و بازیابی اطلاعات، تحلیل اطلاعات و علم‌سنجی، برنامه‌ریزی آموزشی، و مانند آنها کاربرد دارد و در ساخت آنتولوژی (هستان‌شناسی) همچون ابزاری برای وب معنایی، پیش‌نیازی کلیدی و بی‌بدیل است. در اصطلاح‌نامه، واژه‌ها ساختار می‌یابند و مفاهیم با اصطلاح‌ها بیان و به گونه‌ای مرتب می‌شوند که پیوند میان این مفاهیم، آشکار و روشن و اصطلاح‌های مرجح با سرمدخل‌هایی برای مترادف‌ها یا شبه‌مترادف‌ها همراه باشند. اصطلاح‌نامه، راهنمای نمایه‌ساز و کاربر برای گزینش اصطلاح یکسان یا ترکیب یکسانی از اصطلاح‌های مرجح برای بیان یک موضوع است. با ساخت واژه‌های استاندارد برای یک موضوع، یک‌دستی مدخل‌ها در نمایه‌سازی فراهم می‌شود. با تفکیک هم‌نویسه‌ها/هم‌نگاره‌ها به گونه‌ای که هر یک تنها به یک معنی رهنمون شود و گزینش اصطلاح مرجح برای نمایه‌سازی از میان مجموعه‌ای از مترادف‌ها و شبه مترادف‌ها انجام گیرد، کاربر و نمایه‌ساز برای یک مفهوم، اصطلاح همانند یا نزدیک‌تری را بر خواهند گزید.

پردازش زبان طبیعی در پی ساخت و گسترش رویکردهای هوش مصنوعی برای ماشین‌کردن فرایند درک و ساخت یک زبان طبیعی است که به ترجمه، خلاصه‌سازی، درست کردن املا، و مانند آنها به شکل خودکار می‌انجامد؛ این امر با کاربست الگوریتم‌ها و ساختارهای داده در علوم رایانه، به‌ویژه یادگیری ماشین انجام می‌شود. پردازش متن، انجام خودکار بسیاری از پردازش‌ها روی زبان مانند ترجمه، استخراج اطلاعات از متن‌ها، خلاصه‌سازی، استخراج واژه‌های کلیدی، یافتن دستبرد علمی، خوشه‌بندی متن و مدل‌سازی موضوع، نمایه‌سازی، و درست کردن املا، نیاز به مجموعه‌ای از ابزارها برای پیش‌پردازش و آماده‌سازی متن دارد. این ابزارها در دو دسته‌اند که گاهی آمیخته آنها نیز به کار می‌رود. دسته نخست، روش‌های وابسته به زبان هستند که بر پایه برخی از قواعد نحوی و ساختاری زبان انجام می‌شوند. روش‌های دسته دوم مستقل از زبان هستند و بیشتر بر پایه پیکره‌های زبانی و با کاربرد روش‌های یادگیری ماشینی به انجام می‌رسند. طراحی و ساخت این ابزارها برای زبان‌های گوناگون به شیوه‌های گوناگون و ویژه هر زبان انجام می‌شود. ابزارهای پردازش زبان طبیعی برای شناسایی جمله، جداسازی، شناسایی ماهیت اسم‌ها، شبکه واژه‌ها، ریشه‌یاب، شناسایی همانندی، شناسایی گروه‌های نحوی، برجسب‌دهنده نقش نحوی، نشانه‌گذاری در پیکره‌های متنی، تعیین‌کننده هم‌مرجع‌ها، و برجسب‌گذاری اجزای واژگانی کلام به کار می‌روند. کاربردهای مبتنی بر درک زبان طبیعی به واکاوی ژرف‌تری از درون‌داد (متن یا گفتار) نیاز دارد. کاربست تکنیک‌های کلان‌داده و یادگیری ماشین برای درک ماشینی متن‌های خبری و علمی، نمونه‌ای از این کاربردها هستند. ارتباط میان انسان و ماشین و گفت‌و شنود آنها دسته دیگری از کاربردهای پردازش زبان طبیعی هستند. برای

نمونه، خدمات بانکی می‌تواند بر پایه ارتباط میان مشتری و ماشین ارائه شود. طراحی سیستم‌های تبدیل متن به گفتار و گفتار به متن، سیستم‌های ترجمه گفتاری، سیستم‌های گفت‌و شنود، و بازشناسی گفتار از دیگر کاربردهای پردازش زبان طبیعی هستند. انجام بسیاری از طرح‌های پژوهشی در زمینه پردازش زبان طبیعی و پردازش متن نیاز به کاربرست پیکره‌های متنی تخصصی دارد. برای چنین کارهایی بهتر است به جای پیکره‌های متنی عمومی در بازار، پیکره‌هایی بر پایه داده‌های علمی و فنی پایگاه‌های داده ایرانداک ساخته شود. از آنجایی که این داده‌ها از پایان‌نامه‌ها و رساله‌های بروز به دست می‌آیند، می‌توانند در بسیاری از کارهای پژوهشی به کار روند و دستاوردهای بهتری را به بار آورند.

سامانه‌ها

ایرانداک برای پیشبرد اهداف خود در زمینه فناوری اطلاعات، سامانه‌های گوناگونی را فراهم آورده است. که از ابزارهای پردازش زبان طبیعی جهت بازیابی اطلاعات و پردازش متون فارسی استفاده می‌کنند.

سامانه ملی ثبت پایان‌نامه، رساله، و پیشنهاد

هر چند سامانه ملی ثبت پایان‌نامه، رساله، و پیشنهاد از سال ۱۳۸۶ و بر پایه بخش‌نامه وزارت عتف (شماره ۴۳۸۹/۱۲۲۳۸ تاریخ ۱۳۸۶/۸/۲۰) به‌راه افتاد، ولی این سامانه در آبان ۱۳۹۵ برای پشتیبانی از آیین‌نامه پیش‌گفته در نشانی sabt.irandoc.ac.ir درست شد. دانشجویان پس از تصویب پیشنهاد خود، به این سامانه مراجعه و اطلاعات پیشنهاد خود را ثبت و شناسه ره‌گیری دریافت می‌کنند. «پارسا»ها نیز پس از ثبت و بارگذاری فایل تمام‌متن و تأیید ایرانداک، آماده پذیرش دانشگاه می‌شوند. دریافت شناسه ره‌گیری به‌منزله بارگذاری پایان‌نامه یا رساله در سامانه است. سپس دانشگاه وارد سامانه می‌شود و اطلاعات دانشجو را بررسی و اگر تأیید کند، در سامانه ثبت می‌شود. اطلاعات ثبت شده در این سامانه، برای همانندجویی به کار می‌روند و در دسترس همگان نیز گذارده می‌شوند.

ورود به سامانه

رایانامه

گذرواژه

ورود به سامانه

☐ مرا به یاد بسپار گذرواژه‌ام را فراموش کرده‌ام. [ویرایش رایانامه یا شماره ملی](#)

آگهی‌ها

ثبت پیشنهاد، رساله، و پایان‌نامه

از این پس نیازی به فرستادن لوخ فشرده (سی-دی) پایان‌نامه‌ها و رساله‌ها و نسخه چاپی آنها به ایرانداک نیست.

آیین‌نامه اجرایی ثبت پایان‌نامه و رساله در ایرانداک ابلاغ شد.

درباره

آیین‌نامه ثبت و اشاعه پیشنهادها، پایان‌نامه‌ها، و رساله‌های تحصیلات تکمیلی و صیانت از حقوق پدیدآوران در آنها (شماره ۱۹۵۹۳۹/و تاریخ ۱۳۹۵/۹/۴) از همه دانشگاه‌ها، پژوهشگاه‌ها، و مراکز آموزش عالی، پژوهشی، و فناوری دولتی و غیردولتی زیر نظر وزارت علوم، تحقیقات، و فناوری خواسته است که فایل تمام‌متن این مدارک را در پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) ثبت کنند. برای آسان‌سازی این فرایند، ایرانداک سامانه ملی ثبت پایان‌نامه، رساله، و پیشنهاد را در اختیار دانشجویان گرامی قرار داده است...

پایگاه اطلاعات علمی ایران (گنج)

پایگاه اطلاعات علمی ایران (گنج)، دستاورد کاربرد فناوری اطلاعات در خدمات مدیریت اطلاعات علم و فناوری در ایران است. پیشینه این بخش از خدمات ایرانداک به دهه چهل خورشیدی بازمی‌گردد که کار گردآوری، سازمان‌دهی، و اشاعه اسناد علمی و فناوری در مرکز اسناد علمی آغاز شد. بیشترین داده‌های روزآمد گنج از پایان‌نامه‌ها و رساله‌ها و پیشنهادهاست، ولی قلمرو آن گسترده است و طرح‌های پژوهشی، گزارش‌های دولتی، و مقاله‌های علمی را نیز در بر دارد. کار پایگاه اطلاعات علمی ایران، نخستین بار در سال ۱۳۷۰ در رایانه‌های «مین فریم» و محیط نرم‌افزار «سی‌دی‌اس‌آی‌اس‌آی‌اس» شروع شد. در سال ۱۳۷۳ این پایگاه با فناوری «دایال‌آپ» در شبکه‌ای با نام «صبا» از دور در دسترس کاربران قرار گرفت. سال ۱۳۷۸ بود که اینترنت به ایرانداک وارد شد و دسترسی به گنج نیز در نشانی ganj.irandoc.ac.ir در وب فراهم شد. آخرین ویرایش گنج در آذر ۱۳۹۵ رونمایی گردید. این پایگاه با صدها هزار پیشینه از بزرگ‌ترین گنجینه‌های علمی کشور و هم‌اکنون مرجع بسیاری از پژوهشگران ایران و جهان است. این سامانه بنیان سامانه‌های دیگری در ایرانداک همچون سامانه پیشینه پژوهش نیز هست.

The screenshot displays the Ganj website interface. At the top, there is a navigation bar with the logo and name 'پایگاه اطلاعات علمی ایران' (Ganj). Below the navigation bar, there is a search bar with a dropdown menu for 'جستجو' (Search) and a text input field. The main content area shows a list of search results with columns for 'شماره' (Number), 'عنوان' (Title), and 'نوع سند' (Document Type). Below the search results, there are several statistics and links. On the left, there is a section for 'برگزاری دوره آموزشی تخصصی' (Specialized Educational Course). In the center, there is a section for 'سامانه پیشینه پژوهش' (Research Precedence System). On the right, there is a section for 'سامانه همانندجو' (Similarity System).

گنجینه‌ها	شماره مدارک	تمام متن
۹۶۱۳۵۸	۶۳۰۲۲۹	تمام متن

شماره	عنوان	نوع سند
۶۵۱۸۲۵	شماره پایان‌نامه‌های کارشناسی ارشد	تمام متن
۵۶۵۱۵۸	شماره پیشنهادها	تمام متن
۱۸۶۹۹۷	شماره پایان‌نامه‌های کارشناسی ارشد	تمام متن
۹۷۵۳۹	شماره رساله‌های دکتری	تمام متن
۲۵۰۰۷	شماره پیشنهادها	تمام متن
۶۵۰۷۱	شماره رساله‌های دکتری	تمام متن

سامانه پیشینه پژوهش

در پاسخ به خواست جامعه علمی کشور برای هدفمندسازی و افزایش کارایی و اثربخشی

پژوهش‌ها، ایرانداک سامانه پیشینه پژوهش را با پشتوانه در حال افزایش صدها هزار «پارسا» از آبان ۱۳۹۳ راه‌اندازی کرده و در نشانی pishineh.irandoc.ac.ir در اختیار همه پژوهشگران کشور قرار داده است. راه‌اندازی سامانه پیشینه پژوهش، گامی در کمک به پیشبرد پژوهش و گسترش علم و فناوری با آگاهی از پژوهش‌های انجام شده و زمینه‌سازی برای پیش‌گیری از دوباره‌کاری در پژوهش است. در این سامانه، با دریافت عنوان و کلیدواژه‌های پژوهش از کاربر، با جست‌وجوی کارشناسان آموزش‌دیده در پایگاه اطلاعات گنج، اطلاعات پایان‌نامه‌ها و رساله‌های انجام شده به آگاهی کاربر می‌رسد.

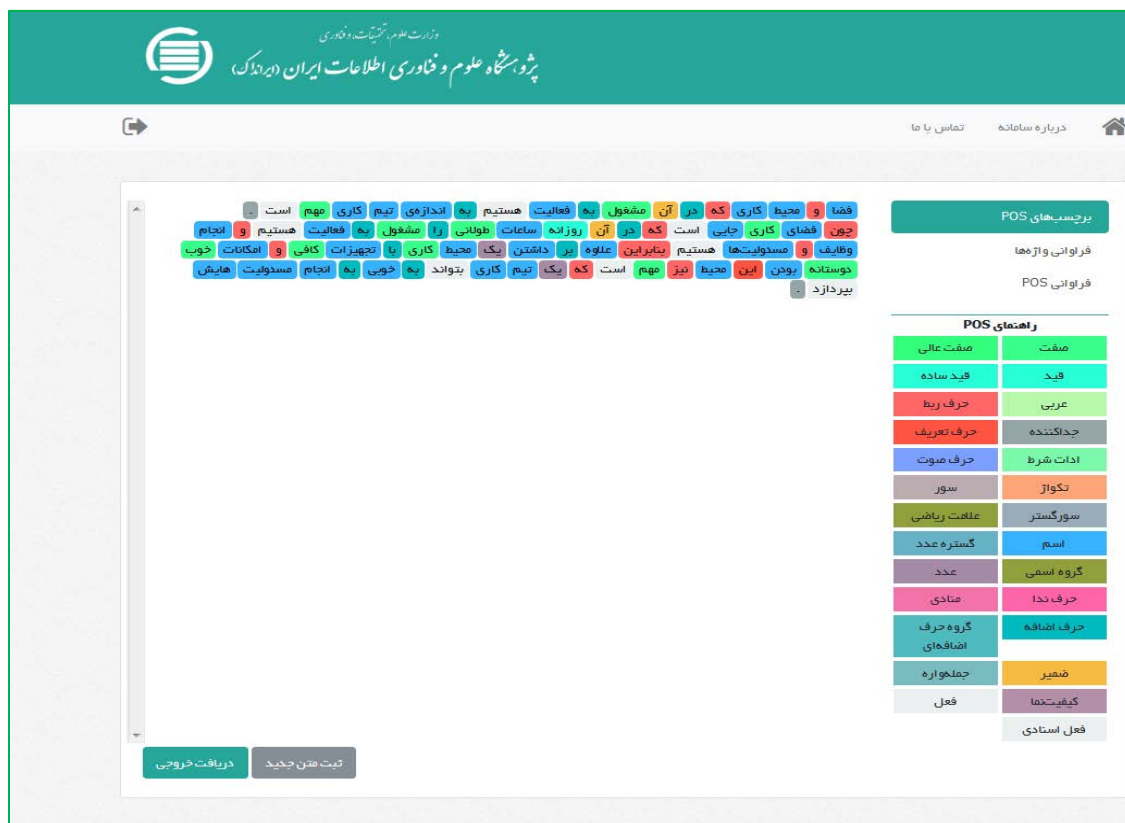
سامانه همانندجو

برای اجرای سیاست‌های کلی علم و فناوری و در پاسخ به خواست جامعه علمی کشور برای گسترش اخلاق پژوهش و پشتیبانی از مالکیت فکری و معنوی و هم‌چنین پیش‌گیری از بدرفتاری‌های علمی، به‌ویژه سایه‌نویسی و دستبرد علمی در نگارش پایان‌نامه‌ها و رساله‌ها، ایرانداک سامانه «همانندجو» را در نشانی tik.irandoc.ac.ir با پشتوانه تمام‌متن و فزاینده ۳۰۰ هزار «پارسا» ۲۰ آذر ۱۳۹۵ راه‌اندازی کرد. همانندجویی در نوشتار پایان‌نامه‌ها و رساله‌ها، گامی در کمک به نگه‌داشت حقوق پدیدآوران و گسترش علم و فناوری و زمینه‌سازی برای دسترسی آزاد همگان به اطلاعات است.



سامانهٔ برچسب‌گذاری اجزای کلام برای متون فارسی (سامانه سبا)

برچسب‌گذاری اجزای کلام (Part Of Speech tagging)، عمل انتساب برچسب به کلمات تشکیل‌دهندهٔ یک متن یا یک پیکره است. این برچسب‌گذاری براساس مقولهٔ آن کلمه در متن (مانند «اسم»، «فعل»، «قید»، «صفت»، و غیره) صورت می‌پذیرد. برچسب‌گذاری اجزای واژگانی کلام، عملی کاربردی در بسیاری از حوزه‌های پیشرفته‌تر پردازش زبان طبیعی از جمله ترجمهٔ ماشینی، خطایاب، تبدیل متن به گفتار، بازیابی اطلاعات، موتورهای جستجو و کمک به مدل‌های آماری است. سامانهٔ برچسب‌گذاری اجزای کلام (که دسترسی به آن در نشانی saba.irandoc.ac.ir فراهم شده است)، امکان وارد کردن متون گوناگون فارسی، برچسب‌گذاری مقوله‌ای کلمات تشکیل‌دهندهٔ متون، مشاهدهٔ فهرست کلمات برچسب‌خورده همراه با فراوانی آن کلمات در متن، مشاهدهٔ فراوانی برچسب‌ها در متن، مشاهدهٔ فهرست اسم‌ها به ترتیب فراوانی آنها در متن، رفع ابهام از برچسب برخی از هم‌نگاره‌های اسمی و صفتی فارسی و دریافت خروجی هر کدام از فهرست‌ها را فراهم می‌آورد. طراحی سامانهٔ برچسب‌گذاری اجزای کلام، مبتنی بر کدهای ابزار برچسب‌گذاری «هضم» است. اما سعی شده است از طریق به کار بردن نرم‌افزار مجهز به رفع ابهام از برچسب نحوی هم‌نگاره‌های اسمی و صفتی مختوم به «ی» در فارسی، که فراوانی بالایی در پیکره‌های متنی فارسی دارند، عملکرد این ابزار بهبود یابد. نرم‌افزار مذکور مبتنی بر الگوهای حساس به بافت نحوی مستخرج از پیکرهٔ «بی‌جن‌خان» است. پس از ارزیابی‌های انجام شده، تنها الگوهایی که تأثیر مثبت در برچسب‌زنی داشته‌اند به برچسب‌زن «هضم» (که مجدداً با «پیکرهٔ بی‌جن‌خان» آموزش دیده است)، اضافه شده است.



سامانه هوشمند استخراج پژوهشگران مشابه

در پایگاه داده‌های ثبت و گنج پژوهشگاه، یک پژوهشگر مشخص به اشکال مختلف ثبت شده است. دلایلی مانند چند املائی بودن برخی نام‌ها، عدم دقت دانشجو، غلط‌های املائی و ... را می‌توان برای علت گوناگونی نام پژوهشگر نام برد. همچنین اقلام اطلاعاتی شناسه ملی، شماره شناسنامه، تاریخ تولد در حجم عظیمی از داده‌ها معتبر نمی‌باشند. سامانه هوشمند استخراج پژوهشگران مشابه با استفاده از الگوریتم‌های زبان‌شناسی رایانشی و تطبیق هوشمند رشته، پژوهشگران یکناپی که به اشکال گوناگون ثبت شده‌اند را استخراج می‌کند.

این سامانه نقش بسزایی در پاکسازی و تجمیع پژوهشگران داشته و امکان طراحی و پیاده‌سازی سرویس‌های ارزش افزوده را میسر می‌سازد. هم‌اکنون حدود ۸۰۰ هزار پژوهشگر در این سامانه در حال تجمیع و پاکسازی می‌باشد. این سامانه جهت استفاده معاونت اطلاعات علم و فناوری ایران در پژوهشگاه طراحی و پیاده‌سازی گردیده است.

IRANDOC

جستجوی هوشمند پژوهشگران

جستجو تقریبی 100 مقدم جرگی 100 نصراله جستجو با نام و نام خانوادگی جستجو با شناسه ملی

به روز رسانی پایگاه داده

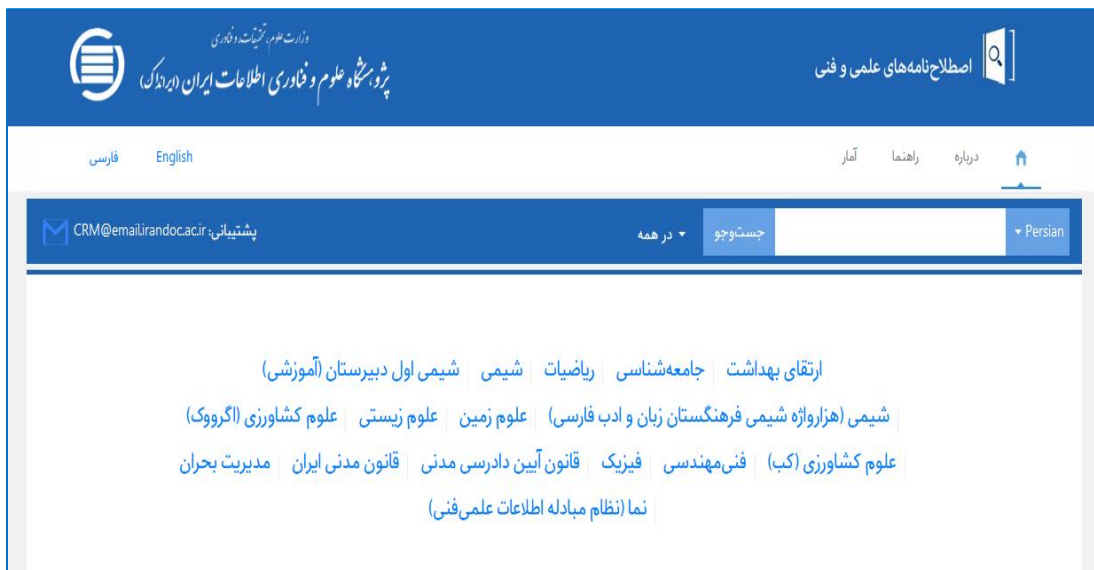
ردیف	شماره پرونده	کد ملی	نام	نام خانوادگی	فنی و مهندسی	علوم انسانی	علوم پایه	هنر و معماری	کشاورزی	علوم پزشکی	ناپزشکی	نام خانوادگی	نام خانوادگی
1	14319	95.0	نصراله	عالم جرگی	49	0	1	0	0	0	0	1	0
2	119772	95.0	نصراله	عالم جرگی	50	0	1	0	0	0	0	0	0
3	128125	95.0	نصراله	عالم جرگی	1	0	0	0	0	0	0	0	0
4	349969	95.0	نصراله	عالم جرگی	7	0	0	0	0	0	0	0	0
5	347202	95.0	نصراله	عالم جرگی	10	0	0	0	0	0	0	0	0
6	422765	95.0	نصراله	عالم جرگی	0	0	2	0	0	0	0	0	0
7	469954	95.0	نصراله	عالم جرگی	3	0	0	0	0	0	0	0	0
8	469548	95.0	نصراله	عالم جرگی	1	0	0	0	0	0	0	0	0
9	675790	95.0	نصراله	عالم جرگی	1	0	0	0	0	0	0	0	0
10	692023	95.0	نصراله	عالم جرگی	1	0	0	0	0	0	0	0	0
11	698555	95.0	نصراله	عالم جرگی	1	0	0	0	0	0	0	0	0

زمان اجرا: 1.51 ثانیه تعداد پروندهها: 798128

سامانه جامع اصطلاح‌نامه‌های علمی و فنی

در این سامانه؛ تمامی اصطلاح‌نامه‌هایی که تا کنون توسط ایرانداک منتشر شده‌اند بارگذاری شده و امکان جستجوی واژه‌ها به صورت جداگانه در هر اصطلاح‌نامه و جستجو در چندین اصطلاح‌نامه به صورت همزمان، وجود دارد. همچنین امکان مشاهده ساختار درختی واژگان به صورت نمایه‌الفبایی به دو زبان فارسی و انگلیسی برای واژه‌ها وجود دارد. اصطلاح‌نامه‌ها گنجینه‌ای از واژه‌ها هستند. میان این واژه‌ها، نظامی شبکه‌ای و مفهومی برقرار است. اصطلاح‌نامه‌ها در سازماندهی اطلاعات، بازیابی از پایگاه‌های اطلاعات، تولید هستان‌شناسی‌ها (آنتولوژی‌ها)، تحلیل اطلاعات، ترجمه ماشینی و نمایه‌سازی ماشینی، و تولید محتوای آموزشی کاربرد دارند. ایرانداک کار روی اصطلاح‌نامه‌ها را از سال ۱۳۵۰ با نگارش دست‌نامه «قواعد و قراردادهای ساختن اصطلاح‌نامه» آغاز کرد و در سال ۱۳۷۵ نخستین اصطلاح‌نامه خود را با نام «نظام مبادله اطلاعات علمی فنی (نما)» به چاپ رساند. در سال‌های ۱۳۷۶ تا ۱۳۸۶ نیز طراحی و نگارش اصطلاح‌نامه دو زبانه فارسی‌انگلیسی فراگیر در حوزه‌های گوناگون علم و فناوری مانند ریاضی، فیزیک، شیمی، علوم زیستی، علوم زمین، مهندسی، کشاورزی، مدیریت بحران، ارتقای بهداشت، و جامعه‌شناسی با بیش از ۱۰۰ هزار واژه انجام شد. در سال ۱۳۹۸، چند ده هزار واژه دیگر به واژه‌های پیشین افزوده شد و روابط معنایی میان اصطلاحات آنها گسترش و بهبود یافت و نسخه دوم سامانه اصطلاح‌نامه‌ها نیز با فارسی‌سازی نرم‌افزار منبع‌باز «اسکاسموس» در دسترس همگان گذارده شد. دستاورد این خدمت ایرانداک هم‌اکنون به شانزده اصطلاح‌نامه رسیده است. این وبگاه می‌تواند

راهنمای مناسبی برای کاربران و نمایه‌سازان در سازماندهی و بازیابی اطلاعات باشد. ایراندوک به عنوان یکی از مراجع ملی در این حوزه، آمادگی دارد تا نسبت به طراحی، تولید، تدوین و روزآمدسازی اصطلاح‌نامه‌ها، تدوین فرا اصطلاح‌نامه‌ها (متاتزاروس‌ها)، و ساخت وب‌سرویس برای اصطلاح‌نامه‌ها اقدام نماید. در تابستان ۱۳۹۸، ده‌ها هزار واژه دیگر به واژه‌های پیشین در اصطلاح‌نامه‌ها افزوده شد و روابط معنایی میان آنها گسترش و بهبود یافت و نسخه دوم سامانه اصطلاح‌نامه‌ها نیز با امکانات بسیار در نشانی esn.irandoc.ac.ir به دست معاون گرامی پژوهش و فناوری وزارت عتف رونمایی شد.



پژوهش

یکی از رویکردهای کلیدی ایرانداک به گسترش خط و زبان فارسی، راهاندازی گروه پژوهشی و انجام پژوهش در این زمینه بوده است.

پژوهشکده، گروه و اعضای هیئت علمی

ایرانداک در تیر ۱۳۹۶ دو گروه پژوهشی را تحت عناوین «زبان‌شناسی رایانشی» و «اصطلاح‌شناسی و هستان‌شناسی» در شورای گسترش آموزش عالی به تصویب رساند که هر دو در پژوهشکده «علوم اطلاعات» قرار گرفته‌اند. این دو گروه هم‌اکنون شش عضو هیئت علمی دارد.



دکتر ملوک السادات حسینی بهشتی
دانش‌آموخته دکتری تخصصی زبان‌شناسی از دانشگاه تهران
پژوهشکده علوم اطلاعات
گروه پژوهشی اصطلاح‌شناسی و هستان‌شناسی



دکتر سیدمهدی سمائی
دانش‌آموخته دکتری تخصصی زبان‌شناسی از دانشگاه تهران
پژوهشکده علوم اطلاعات
گروه پژوهشی اصطلاح‌شناسی و هستان‌شناسی



دکتر الهام علایی ابوذر
دانش‌آموخته دکتری تخصصی زبان‌شناسی همگانی از
دانشگاه تهران
پژوهشکده علوم اطلاعات
گروه پژوهشی زبان‌شناسی رایانشی



تقی رجیبی
دانش‌آموخته کارشناسی ارشد شیمی فیزیک از
دانشگاه تربیت مدرس
پژوهشکده علوم اطلاعات
گروه پژوهشی اصطلاح‌شناسی و هستان‌شناسی



دکتر نصراله پاکنیت
دانش آموخته دکتری تخصصی ریاضی (علوم کامپیوتر) از
دانشگاه شهید بهشتی
پژوهشکده علوم اطلاعات
گروه پژوهشی زبان‌شناسی رایانشی

طرح‌های پژوهشی

اصطلاح‌نامه‌نگاری

اصطلاح‌نامه گنجینه‌ای از واژه‌هاست که افزون بر نظم الفبایی موجود در فرهنگ‌ها، دارای نظامی شبکه‌ای و مفهومی میان واژه‌های یک یا چند حوزه از دانش بشری است. این نظام مفهومی، ارتباطات معنایی گوناگون میان واژه‌ها را نشان می‌دهد. اصطلاح‌نامه‌ها در سازماندهی اطلاعات برای ذخیره و نمایه‌سازی، بازیابی اطلاعات به کمک شبکه معنایی، تحلیل اطلاعات و علم‌سنجی، مصورسازی پژوهش، تحلیل متن و پردازش زبان علمی، طراحی نقشه‌های مفهومی علم در حوزه‌های گوناگون، طراحی ساختار و محتوای آموزشی و آسان‌سازی یاددهی و یادگیری، داده‌کاوی، ترجمه ماشینی و نمایه‌سازی ماشینی، معادل‌یابی فارسی واژه‌های بیگانه و شناخت اصطلاح‌های استاندارد به کار می‌روند. اصطلاح‌نامه در ساخت آنتولوژی (هستان‌شناسی) به مثابه ابزاری برای وب معنایی، نیازی کلیدی و بدون جایگزین است. ایرانداک در سال ۱۳۷۶ نگارش نخستین اصطلاح‌نامه خود را با نام اصطلاح‌نامه فنی و مهندسی آغاز و تا سال ۱۳۸۶، ده اصطلاح‌نامه دو زبانه با نام‌های؛ ارتقای بهداشت، جامعه‌شناسی، ریاضیات، شیمی، علوم زمین، علوم زیستی، علوم کشاورزی، فنی و مهندسی، فیزیک و مدیریت بحران؛ را با بیش از یکصد هزار واژه، تدوین و به جامعه علمی کشور تقدیم کرده است. ایرانداک با سازمان‌های جهانی اصطلاح‌شناسی مانند «ترم نت» و «اینفوترم» نیز همکاری دارد و توانسته است با کمک یونسکو برخی از اصطلاح‌نامه‌های خود را در سایت‌های جهانی نیز بگذارد.

طراحی نقشه مفهومی برای اصطلاح‌نامه زبان‌شناسی

در زبان فارسی اصطلاح‌نامه دوزبانه فارسی-انگلیسی زبان‌شناسی موجود نیست. گام اول برای تدوین اصطلاح‌نامه زبان‌شناسی دوزبانه فارسی-انگلیسی، مشخص نمودن تقسیم‌بندی دقیق علم زبان‌شناسی به صورت نظری است. در این پروژه نقشه مفهومی فارسی برای اصطلاح‌نامه دو زبانه فارسی-انگلیسی زبان‌شناسی طراحی و تدوین شده است تا در پروژه‌ای دیگر بتوان بر مبنای آن، اصطلاح‌نامه زبان‌شناسی دو زبانه فارسی-انگلیسی را تدوین کرد. حدود ۱۰۰۰ واژه برای طراحی نقشه مفهومی زبان‌شناسی در نظر گرفته شده است. سرفصل دروس و کتب معتبر دانشگاهی در مقاطع تحصیلات تکمیلی زبان‌شناسی و منابع مرجع و فرهنگ‌های تخصصی زبان‌شناسی انگلیسی و فارسی، در گرایش‌های مختلف آواشناسی، واج‌شناسی، صرف، نحو و معنی‌شناسی مورد استفاده قرار گرفته‌اند. همچنین رویکردهای تاریخی علم زبان‌شناسی از منظر سنتی، ساختاری و شناختی مورد توجه خواهند بود. در این پروژه برای دیداری‌سازی نقشه‌های طراحی شده از نرم‌افزار Visio استفاده شده است.

ایجاد هستی‌شناسی از اصطلاح‌نامه شیمی ایرانداک

هستی‌شناسی، واژگان و مفاهیم مشترک مورد استفاده برای توصیف و ارائه یک حوزه از دانش را معین می‌کند و می‌تواند مبنایی برای سازماندهی و مدیریت دانش در یک حوزه خاص باشد. با این وجود، ایجاد هستی‌شناسی بسیار زمان‌بر است. به منظور تسهیل در این امر می‌توان از منابع دیگر دانش مانند واژه‌نامه‌ها و اصطلاح‌نامه‌ها استفاده کرد. اصطلاح‌نامه‌ها با توجه به دارا بودن اطلاعات معنایی و ساختار سلسله‌مراتبی مفاهیم و نیز روابط معنایی میان آن‌ها، منابع مناسبی برای ساخت هستی‌شناسی هستند. در این پژوهش، اصطلاح‌نامه شیمی که پیش‌تر در پژوهشگاه علوم و فناوری اطلاعات تدوین شده است، به عنوان مبنای ساخت هستی‌شناسی شیمی مورد استفاده قرار گرفت. طراحی مفهومی هستی‌شناسی با استفاده از روش «مت‌آنتولوژی» و بر اساس مفاهیم و روابط موجود در اصطلاح‌نامه صورت گرفت. در واقع این اطلاعات از محیط اصطلاح‌نامه‌ای به نرم‌افزار پروتزه انتقال داده شد و سپس روابط معنایی آن گسترش یافته و غنی‌سازی صورت گرفت. در این طرح ۳۰۰۰ مفهوم اصطلاح‌نامه شیمی به صورت روابط معنایی هستی‌شناسی شیمی تبدیل و غنی‌سازی شد.

مفهوم‌سازی و روش تدوین فرااصطلاح‌نامه

فرااصطلاح‌نامه پایگاه اطلاعات واژگانی جامع، بزرگ و گسترش‌یابنده چندمنظوره و چندزبانی حاصل از انطباق، اتصال، ادغام و تلفیق اصطلاح‌نامه‌ها و زبان‌های مقید موجود است که در اصل به قصد استفاده توسعه‌دهندگان آن تهیه و تدوین می‌شود. فرااصطلاح‌نامه می‌تواند اطلاعات مورد نیاز برنامه‌های رایانه‌ای را برای ایجاد داده‌های استاندارد تأمین نماید، پرسش‌های کاربران را تفسیر و از طریق پالایش و جرح و تعدیل سؤالات با آن‌ها تعامل داشته باشد و اصطلاحات کاربران را به واژگان مورد استفاده در منابع مرتبط برگرداند. فرااصطلاح‌نامه معمولاً در حوزه معینی از دانش شکل می‌گیرد. در ایران‌داک اصطلاح‌نامه‌های تخصصی در دهه ۱۳۸۰ تدوین شده است و در سال‌های ۹۷-۱۳۹۶ روزآمدسازی شده‌اند. ایران‌داک، علاوه بر اصطلاح‌نامه‌ها، مجموعه‌ای از واژه‌نامه‌ها و واژگان تخصصی را نیز تدوین و به چاپ رسانده است. برای جستجو و کاوش دقیق‌تر ضرورت دارد که این واژگان و اصطلاح‌نامه‌ها با یکدیگر تلفیق یابند تا از ریزش کاذب در جستجوی اطلاعات ممانعت به عمل آید و سرعت دستیابی هدفمند به اطلاعات تخصصی فراهم گردد. در این پژوهش به بررسی منابع گوناگون با موضوع فرااصطلاح‌نامه و شناسایی مفاهیم وابسته به آن، بیان کاربرد و نیز مراحل تدوین آن پرداخته شده است.

تکمیل، توسعه و دوبانه‌سازی اصطلاح‌نامه آیین دادرسی مدنی ایران

در این پژوهش اصطلاح‌نامه آیین دادرسی مدنی ایران بر اساس قانون آیین دادرسی مدنی دادگاه‌های عمومی و انقلاب مصوب ۱۳۷۹ طراحی و تدوین شده و برای هر اصطلاح، برابر انگلیسی آن نیز درج شده است. برای سازماندهی این اصطلاح‌نامه از نرم‌افزار «تزاروس بیلدر» بهره‌برداری شد. مخاطبان این اصطلاح‌نامه، افزون بر کتابداران، فرهنگ‌نویسان، اصطلاح‌نامه‌سازان و حقوقدانان، همه مترجمان رسمی و دیگر مترجمانی هستند که با متون حقوقی سروکار دارند.

اصطلاح‌نامه آیین دادرسی مدنی ایران (بر اساس قانون آیین دادرسی مدنی دادگاه‌های عمومی و انقلاب مصوب ۱۳۷۹)

اصطلاح‌نامه آیین دادرسی مدنی ایران بر اساس قانون آیین دادرسی مدنی دادگاه‌های عمومی و انقلاب مصوب ۱۳۷۹ و با استفاده از نرم‌افزار «تزاروس بیلدر» ایجاد شده است. این اصطلاح‌نامه دربردارنده ۲۳۱۱ اصطلاح مرجح و ۱۵۰ اصطلاح وابسته، ۲۲۸۰ اصطلاح غیرمرجح، و ۵۹ یادداشت دامنه است. شمار ۱۲۲ مورد از این اصطلاح‌ها در سطح نخست، ۱۲۲۹ مورد در سطح دوم، ۶۹۱ مورد در سطح سوم، ۲۱۳ مورد در سطح چهارم، ۳۳ مورد در سطح پنجم، نه مورد در سطح ششم، و سه مورد

در سطح هفتم هستند.

توسعه و روزآمدسازی اصطلاح‌نامه‌های ایرانداک

توسعه علم و فناوری، تولید واژه‌های جدید، افزایش انتشارات علمی، واژه‌هایی که توسط کاربران وارد نظام سازماندهی اطلاعات می‌شود، و تغییراتی که در روابط بین مفاهیم و موضوعات ایجاد می‌شود از مجموعه دلایل روزآمدسازی اصطلاح‌نامه‌ها هستند. در این پژوهش، افزایش کمی، ویرایش و توسعه روابط معنایی میان واژه‌های اصطلاح‌نامه‌های علمی و فنی تدوین شده در پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)، صورت گرفته است. فرایند روزآمدسازی و توسعه، شامل مراحل مرور ادبیات و تهیه منابع، طراحی و بررسی نقشه‌های مفهومی، ترجمه واژگان، تعیین روابط معنایی، یکسان‌سازی، واپایش روابط و خطایابی، تولید خروجی‌های اصطلاح‌نامه‌ای، ویرایش، و انتشار است. در سال ۱۳۹۸، چند ده هزار واژه دیگر به واژه‌های پیشین افزوده شد و روابط معنایی میان اصطلاحات آنها گسترش و بهبود یافت و نسخه دوم سامانه اصطلاح‌نامه‌ها نیز با فارسی‌سازی نرم‌افزار منبع‌باز «اسکاسموس» در دسترس همگان گذارده شد. دستاورد این خدمت ایرانداک هم‌اکنون به شانزده اصطلاح‌نامه رسیده است. این وبگاه می‌تواند راهنمای مناسبی برای کاربران و نمایه‌سازان در سازماندهی و بازیابی اطلاعات باشد.

توسعه و روزآمدسازی اصطلاح‌نامه شیمی

روزآمدسازی و ویرایش مجموعه واژگان اصطلاح‌نامه شیمی تدوین شده در ایرانداک موضوع این پژوهش است. در روزآمدسازی، اکثر مراحل که در تألیف و تدوین و یا ترجمه اصطلاح‌نامه وجود دارد طی خواهد شد. این مراحل در سطح کلان عبارتند از: ۱- مرور ادبیات و تهیه منابع؛ ۲- طراحی و یا بررسی شبکه مفهومی؛ ۳- ترجمه واژگان؛ ۴- تعیین روابط معنایی؛ ۵- یکسان‌سازی؛ ۶- کنترل روابط و خطایابی؛ ۷- تولید خروجی‌های اصطلاح‌نامه‌ای؛ ۸- ویرایش؛ ۹- چاپ و انتشار. نمایه‌سازان در بسیاری از موارد هنگام استفاده از اصطلاح‌نامه در سازماندهی اطلاعات، با اصطلاحاتی مواجه می‌شوند که آنها را در اصطلاح‌نامه‌های موجود نمی‌یابند. رشد علم و فناوری و افزایش انتشارات علمی و انتخاب توصیفگرهای مناسب‌تر هنگام سازماندهی اطلاعات از یک سو و نیز دیدگاه‌های متفاوت کاربران نهایی نظام‌های بازیابی اطلاعات از سویی دیگر، روزآمدسازی و گسترش روابط معنایی در اصطلاح‌نامه را ضروری می‌نماید. روزآمدسازی اصطلاح‌نامه شیمی با استفاده از واژگان کنترل شده سامانه پایگاه اطلاعات ایرانداک (گنج) و مجموعه واژگان شیمی مرکز نشر دانشگاهی و نیز مجموعه واژه‌های شیمی

مصوب فرهنگستان زبان و ادب فارسی انجام شد.

طراحی ساختار درختی و تدوین اصطلاح‌نامه شیمی از واژگان مصوب فرهنگستان

واژه‌های فارسی مصوب فرهنگستان زبان و ادب فارسی، فاقد ساختار درختی و تزاروسی است. در این پروژه، این واژه‌ها در یک ساختار اصطلاح‌نامه‌ای ارائه شد. عرضه الگویی به فرهنگستان برای استفاده از شبکه معنایی جهت سازماندهی واژه‌های مصوب، در کلیه گرایش‌ها یکی از اهداف این پژوهش است. واژگان شیمی مصوب فرهنگستان از طریق سایت فرهنگستان و نیز از کتاب "هزار واژه شیمی" فرهنگستان قابل دستیابی است. در طراحی ساختار درختی از روش استقرایی (حرکت از جز به کل) استفاده شد. در این مرحله اشراف موضوعی مهمترین عامل در چینش واژگان بود. بنابراین در حوزه‌های مختلف شیمی از مشاوره متخصصین بهره‌مند شده‌ایم. در تدوین اصطلاح‌نامه و تولید خروجی‌های اصطلاح‌نامه‌ای از نرم‌افزار "تزاروس بیلدر" استفاده شد و تمام روش‌های کلاسیک تدوین اصطلاح‌نامه مورد استفاده قرار گرفت. تولید اصطلاح‌نامه‌ها به دلیل کاربرد مستقیم آن در استاندارد کردن و بازیابی اطلاعات در راستای برنامه استراتژیک و نیز اساس نامه ایرانداک است.

اصطلاح‌نامه دبیرستان و آموزش نوین (اصطلاح‌نامه آموزشی)

در اکثر پژوهش‌های انجام شده در ارتباط با روان‌شناسی یادگیری و آموزش، الگوها و روش‌های تدریس، برنامه‌ریزی آموزشی، آموزش استراتژیک تکنولوژی آموزشی در کنار صدها اصطلاح و مبحث تخصصی همواره با اصطلاح مواد درسی (ماده درسی) و یا محتوای آموزشی برخورد می‌کنیم. مواد یا محتوای آموزشی در شیوه‌های سنتی و مدرن در مبحث تدریس و یادگیری انواع و اقسامی دارند ولی شاکله همه آنها متون کتابی است. ارائه مواد درسی (محتوای آموزشی) در قالب اصطلاح‌نامه تصویری در کنار (یا به جای) متون کتابی به منظور افزایش یادگیری فراگیران فرضیه‌ای است که در این طرح مورد آزمون قرار می‌گیرد. برای اصطلاح‌نامه (تزاروس) دو کاربرد اصلی: ۱. استاندارد کردن اطلاعات؛ ۲. بازیابی اطلاعات ذکر می‌کنند. در این طرح می‌خواهیم نشان دهیم که اول می‌توان محتوای آموزشی (ماده درسی) را در قالب ساختاری اصطلاح‌نامه‌ای ارائه کرد، دوم: این قالب ساختاری به عنوان یک ابزار موثر آموزشی و یا کمک آموزشی در آموزش علوم کاربرد دارد و موجب افزایش یادگیری فراگیران خواهد گردید. در این طرح در ابتدا به عنوان مطالعه موردی کتاب شیمی سال اول دبیرستان انتخاب می‌شود و کل مفاهیم و موضوعاتی که در این طرح مطرح شده‌اند، در قالب یک اصطلاح‌نامه تصویری ارائه خواهد شد. پس با استفاده از دو روش آزمون و پیمایش به سنجش نمرات و نظرات گروه‌های آزمون و کنترل در چند دبیرستان تهران اقدام خواهیم نمود. این طرح از هر دو منظر بنیادی و

کاربردی می‌تواند از مصادیق رابط بین علوم اطلاع‌رسانی، ارتباطات، مهندسی و مدیریت دانش، روان‌شناسی، تکنولوژی آموزشی و نیز رایانه قلمداد گردد.

تدوین روش‌شناسی تحلیل اطلاعات علمی شیمی (علوم پایه) و ترسیم نقشه علمی شیمی در ایران بر اساس آن (با بررسی پایگاه اطلاعاتی پایان‌نامه‌های پژوهشگاه)

در این پژوهش ابتدا الگویی جهت تحلیل اطلاعات علمی پایان‌نامه‌های شیمی (علوم پایه) ارایه شد، سپس با استفاده از این الگو به تحلیل پایان‌نامه‌های شیمی موجود در پایگاه اطلاعات پژوهشگاه علوم و فناوری اطلاعات ایران در یک دهه اخیر پرداخته شد. با استفاده از منابع معتبر شیمی در سطح کارشناسی و تحصیلات تکمیلی و نیز اصطلاح‌نامه شیمی و با تأیید نظر متخصصان در هریک از گرایش‌های چهارگانه شیمی علوم پایه (شیمی فیزیک، شیمی آلی، شیمی معدنی، و شیمی تجزیه) به تهیه نقشه علم شیمی همت گماشتیم. در حقیقت تحلیل اطلاعات علمی شیمی، بدون در اختیار داشتن الگوی مدونی از علم شیمی بسیار مشکل خواهد بود. جست‌وجوی فراوان برای یافتن چنین الگویی بی‌نتیجه بود، لذا الگوی مدونی از علم شیمی (علوم پایه) برای اولین بار به این صورت ارایه شد، که علاوه بر استفاده در تحلیل اطلاعات علمی، قابلیت استفاده آموزشی هم دارد و البته با کمک متخصصان می‌توان به بسط و غنای علمی آن افزود. تحلیل علمی اطلاعات پایان‌نامه‌های شیمی که یک تحلیل توصیفی است می‌تواند در سیاست‌گذاری‌های تحقیقاتی نقشی کلیدی ایفا نماید.

بررسی مفهوم و ابعاد خلاصه همگان‌فهم

ترویج علم، شرایطی را فراهم می‌سازد که گذشته از نخبگان و دانشمندان حوزه‌های مختلف علوم، سایر افراد جامعه نیز درک و بینشی از علم به دست آورند. یکی از ابزارهای ترویج علم، تنظیم و انتشار خلاصه همگان‌فهم است. خلاصه همگان‌فهم، خلاصه کوتاهی از پیشنهادیه یا طرح پژوهشی است که غیر از پژوهشگران و متخصصان همان حوزه، برای عموم به نگارش درآمده است. این خلاصه باید به زبان ساده نوشته شود، از کاربرد واژگان تخصصی دوری شود و هر واژه فنی موجود در آن به خوبی تشریح شود. هدف این پژوهش تشریح مفهوم خلاصه همگان‌فهم، شباهت‌ها و تفاوت‌های مفاهیم همانند با آن، سازمان‌های فعال، الگوهای پیشنهادی، زمینه‌های موضوعی، و اشکال بکارگیری آن در تولیدات علمی بوده است. شباهت‌ها و تفاوت‌های مفاهیم مورد بررسی شامل «چکیده»، «خلاصه»، «خلاصه اجرایی» و «خلاصه همگان‌فهم» از چهار جنبه «مخاطب یا مخاطبان»، «ساختار»، «هدف» و «کاربرد در نوع مدرک»، به صورت مقایسه‌ای ارائه شد.

استخراج فراداده در پارساها با رویکرد الگوریتم‌های دسته‌بندی

فراداده داده‌ای است که اطلاعاتی درباره داده‌ای دیگر ارائه می‌دهد. عنوان، نام دانشگاه، رشته تحصیلی، نام دانشجو، استاد راهنما، چکیده و کلمات کلیدی را می‌توان از مهم‌ترین فراداده‌های پایان‌نامه‌ها و رساله‌ها (پارساها) نام برد. از آنجایی که فراداده‌ها در بازیابی اطلاعات نقش منحصر به فردی را ایفا می‌نمایند، استخراج صحیح و ذخیره آن‌ها در سامانه‌های بازیابی اطلاعات حائز اهمیت است. با توجه به حجم بالای متون علمی و سرعت پایین کنترل دستی فراداده در سامانه ثبت پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)، همواره نمایه‌سازی پارساهای دریافتی عقب‌تر از پارساهای دریافتی بوده و در نتیجه نمایش اطلاعات در سامانه گنج ایرانداک به روز نیست. با توجه به این مساله، افزایش سرعت کنترل فراداده‌ها از موارد اولویت‌دار سامانه ثبت است. برای حل این مساله، در این پژوهش الگوریتمی طراحی شده که قادر است فراداده‌ها را از تمام متن پارساها استخراج نماید. الگوریتم طراحی شده مبتنی بر روش‌های هوشمند دسته‌بندی چند کلاسه بوده و در آن از الگوریتم‌های فرامکاشفه‌ای برای مکاشفه بهتر و دستیابی به مقادیر بهینه دسته‌بند استفاده شده است. ارزیابی روش پیشنهادی نشان می‌دهد در بسیاری از اقلام اطلاعاتی فراداده استخراج هوشمند و خودکار کارا بوده و امکان عملیاتی شدن ایده وجود دارد.

تهیه مجموعه داده استاندارد فارسی برای مساله استخراج خودکار کلیدواژه از اسناد علمی

با توجه به حجم داده‌ها و مستندات ایجاد شده در پایگاه‌های تخصصی علمی نیاز است که روش‌های هوشمند و خودکار برای افزایش دقت و سرعت تخصیص کلیدواژه به این مستندات به کار گرفته شود. توسعه و بکارگیری این روش‌ها نیازمند وجود معیارهای استاندارد برای ارزیابی و بهبود آنهاست. یکی از این معیارها، وجود یک مجموعه داده استاندارد است که حاوی مجموعه‌ای از اسناد و کلیدواژه‌های تخصیص داده شده درست و کامل به آنها باشد. با توجه به عدم وجود مجموعه داده‌ای استاندارد جهت استخراج کلیدواژه از متون فارسی علمی، در این پژوهش، مجموعه داده‌ای از پایان‌نامه‌ها و رساله‌های فارسی به همراه کلیدواژه‌های مرتبط با آنها ایجاد شده است. در این راستا، از مجموعه‌ای از اسناد پایگاه اطلاعات علمی ایران (گنج) که کلیدواژه‌های تخصصی نمایه‌ساز و نویسنده را دارند استفاده شده و در ادامه، براساس فرایندی ماشینی مجموعه کلیدواژه وسیع‌تری، به عنوان کلیدواژه‌های کاندید برای هر سند در نظر گرفته شده و در انتها نیز متخصصین نمایه‌ساز، با بررسی مجدد هر سند و کلیدواژه‌های کاندید، کلیدواژه‌های تخصصی مناسب متناظر با هر سند را انتخاب کرده‌اند. در نهایت ۲۸۳۸ داده در چهار حوزه موضوعی فنی و مهندسی، علوم انسانی، علوم و علوم کاربردی، و هنر در مجموعه داده قرار

گرفته‌اند که هر یک به طور متوسط دارای ۱۰ کلیدواژه هستند. در انتها نیز پیشنهادهایی برای چگونگی بهره‌برداری از مجموعه داده استاندارد توسط جامعه علمی و انتشار آن ارائه شده است.

شناسایی ویژگی‌های نحوی زبان علم در مدارک علمی فارسی

زبان علم یکی از گونه‌های اجتماعی زبان است که طبقات تحصیل کرده در مجامع و بافت‌های علمی آن را به کار می‌بندند. هدف از انجام این پژوهش عرضه تصویری کلی از ویژگی‌های نحوی زبان علم فارسی بوده است. بدین معنی که انواع جملات، انواع وجه، زمان‌های افعال و ساخت‌های نحوی از متون مورد بررسی استخراج و آمار آنها عرضه و محتوای آنها تحلیل شده است. برای شناسایی ویژگی‌های نحوی زبان علم پیرامون ۹۰ مقاله در نشریه‌های علمی-پژوهشی در دو حوزه علمی (علوم انسانی و علوم پایه) و سه رشته دانشگاهی (زبان‌شناسی، زمین‌شناسی، و شیمی) به شکل هدف‌مند نمونه‌گیری شدند. این مقاله‌ها به شکل متن وارد نرم‌افزار «کیو.دی.ای. ماینر لایت»، مطالعه، و کدگذاری شدند. روی هم، پیرامون ۱۴۰۰ کد از این آثار استخراج و تحلیل شدند. یافته‌های پژوهش نشان دادند که از ویژگی‌های نحوی غالب مشترک زبان علم فارسی استفاده از وجه اخباری و زمان حال ساده و ماضی نقلی و گذشته ساده و جملات ساده و وجه مجهول بوده است. وجه اخباری با ۳۷۷ کد و زمان حال ساده با ۳۲۵ کد بیشترین فراوانی را در فرایند کدگذاری به خود اختصاص دادند.

ساخت پیکره متنی از مقاله‌های پژوهش‌نامه پردازش و مدیریت اطلاعات

انجام بسیاری از پژوهش‌های زبان‌شناسی و تصمیم‌گیری‌ها در برنامه‌ریزی‌های زبانی، تنها با استفاده از یک پیکره زبانی امکان‌پذیر است. در پژوهش حاضر پیکره‌ای ساخته شده است که شامل متون مقاله‌های موجود در پژوهش‌نامه پردازش و مدیریت اطلاعات است. ۶۷۷ مقاله و ۴۷۷۴۷۹۰ واژه. موضوع این مقاله‌ها شامل کتابداری و اطلاع‌رسانی، علم اطلاعات و دانش‌شناسی، فناوری اطلاعات، مدیریت اطلاعات، مدیریت دانش، زبان‌شناسی، زبان‌شناسی رایانشی، اصطلاح‌شناسی، هستان‌شناسی و سایر حوزه‌های مرتبط با پردازش اطلاعات است. بنابراین، محتوای این پیکره عمومی نیست و حاوی متون بسیار تخصصی و میان‌رشته‌ای است و از این نظر برای پردازش‌هایی که مستلزم بهره‌گرفتن از متون تخصصی است، بسیار ارزشمند است. پس از نمونه‌گیری و وارد کردن متون مقاله‌ها در پیکره، فراداده مربوط به مقاله‌ها (مانند عنوان مقاله، نام نویسنده/نویسندگان و موضوع مقاله) نیز از طریق کاربرگ ثبت مقالات وارد پیکره شد. سپس، ابتدا نرمال‌سازی ماشینی و به دنبال آن، برچسب‌گذاری ماشینی) نوعاً برچسب‌گذاری اجزای واژگانی کلام (POS) انجام شد و اطلاعات در پایگاه داده ذخیره شده است و در نهایت، برچسب‌ها دستی کنترل شده‌اند. همچنین در این طرح پژوهشی، در کنار ساخت پیکره،

سامانه‌ای طراحی شد به نام «سامانه پیکره‌های ایرانداک» که پیکره متنی ایرانداک در این سامانه قرار گرفت و در آینده می‌توان پیکره‌های دیگری را نیز در این سامانه قرار داد.

راهنمای نگارش خلاصه همگان‌فهم در حوزه مدیریت آب

هدف کلیدی این پژوهش تدوین راهنمایی عملیاتی برای نگارش خلاصه همگان‌فهم است که دربردارنده آموزه‌های کاربردی باشد تا پژوهشگران و یا نهادهای درگیر پژوهش، بر پایه آن بتوانند خود، خلاصه‌هایی همگان‌فهم به رشته نگارش درآورند. رویکرد این پژوهش ترکیبی و در سه گام پیش رفته است. پس از تدوین راهنمای عملیاتی با کمک روش مرور نظام‌مند نوشته‌ها، برای نگارش خلاصه همگان‌فهم، این راهنما توسط متخصصان ترویج علم اعتباریابی شده است. سپس، برای ۱۰ اثر پژوهشی در زمینه مدیریت آب توسط متخصصان این حوزه در یک گروه کانونی خلاصه همگان‌فهم بر پایه راهنمای پیشنهادی تهیه شده است. سرانجام، خلاصه‌های تهیه شده به ۳۵۶ خواننده عامه داده شده است تا درک آنها از خلاصه‌ها و در نتیجه کارآمدی راهنمای تهیه شده سنجش شود. یافته‌های پژوهش نشان دادند که راهنمای نگارش خلاصه همگان‌فهم دست کم دربردارنده سه بُعد و ۲۳ مولفه است. بُعد «شکلی و صوری» دربردارنده مولفه‌های از زبان اول شخص، بهره‌گیری از جمله معلوم، گزینش طرح منطقی، فضای سفید میان سطرها، نیاوردن استناد، پرهیز از کوتاه‌نوشت، بهره‌گیری از جمله‌های کوتاه، محدودیت شمار واژگان، داستان‌گویی، و بهره‌گیری از جمله‌های مثبت؛ بُعد «محتوایی» دربردارنده مولفه‌های پاسخ به چه کسی، چی، کجا، چه زمانی، چگونه، و چرا، توضیح تاثیر پژوهش، نیاوردن واژه تخصصی، پیچیده، و بلند، توجه به ویژگی گروه هدف، بهره‌گیری از سرنویس‌های ساده، بکارگیری بافت، ارائه دلیل انجام پژوهش، و آوردن تصویر/ نمودار/ داده‌نما؛ و بُعد «سنجش و ارزیابی» نیز دربردارنده کنترل میزان خوانش، بازخورد دیگر پژوهشگران، و بازخورد گروه هدف است. نتیجه آزمون «تی مستقل» نشان داد که میان فهم «همگان» از چکیده‌های فنی و خلاصه‌های همگان‌فهم تفاوت معنادار هست و آنان که خلاصه همگان‌فهم را خوانده فهم بهتری از پژوهش‌ها داشته‌اند.

شناسایی ویژگی‌های زبان علم در مدارک علمی فارسی

زبان فارسی نیز همانند دیگر زبان‌ها به گونه‌های اجتماعی متفاوتی تقسیم می‌شود. این گونه‌ها تابع متغیرهای سن و جنسیت و حرفه و جایگاه اجتماعی گویشوران در جامعه زبانی است. زبان علم یکی از گونه‌های اجتماعی زبان فارسی است که طبقات تحصیل کرده در مجامع و بافت‌های علمی آن را به کار می‌بندند. در این پژوهش ویژگی‌های صوری و زبانی بخشی از مدارک علمی فارسی در رشته‌های زبان‌شناسی و زمین‌شناسی و شیمی استخراج شده است. هدف از انجام این پژوهش عرضه تصویری

کلی از ویژگی‌های زبان علم فارسی بوده است. بدین معنی که مقولات واژگانی نظیر اسم و صفت و قید و حروف اضافه و برخی ساخت‌های نحوی نظیر ساخت مجهول و زمان افعال از متون مورد بررسی استخراج و آمار آنها عرضه و محتوای آنها تحلیل شده است. از ویژگی‌های مشترک زبان علم فارسی استفاده از الفاظ فخیم فارسی و عربی در مقوله‌های اسم و صفت و قید و حروف اضافه و کاربرد فراوان ساخت مجهول و استفاده بیشتر از زمان افعال بوده است.

طراحی روشی هوشمند برای دسته‌بندی نوشتارهای علمی فارسی

پایان‌نامه و رساله‌ها (پارساها)یی که در پایگاه داده ثبت پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) ثبت می‌گردد دارای فراداده موضوع اصلی مانند علوم انسانی، فنی و مهندسی، علوم پایه، هنر و معماری، علوم پزشکی و غیره هستند. استخراج هوشمند موضوع اصلی پارساها بوسیله الگوریتم‌های یادگیری ماشین باعث افزایش سرعت فرایند نمایه‌سازی و افزایش کیفیت داده‌های پایگاه اطلاعات علمی ایران (گنج) شود. در این طرح پژوهشی هدف دسته‌بندی هوشمند محتوایی متون به یکی از ۵ دسته موضوع اصلی علوم انسانی، فنی و مهندسی، علوم پایه، هنر و معماری، و علوم پزشکی است. در این پژوهش ابتدا پیش‌پردازش معمول در پردازش زبان طبیعی انجام شده و ویژگی‌های متمایز کننده استخراج شده است. در ادامه، با استفاده از محتوای مجموعه‌ای از پارساها، دسته‌بندی‌های مختلف جهت دسته‌بندی پارساها به دسته‌های موردنظر آموزش داده شد. نتایج ارزیابی‌ها بر روی متون پایان‌نامه و رساله‌های فارسی نشان می‌دهد این روش با ارزیابی با متون بیشتر می‌تواند با دقت و سرعت خیلی خوبی موضوع اصلی پارساها را مشخص کند.

شناسایی مسائل و ارائه راهکارهای زبان‌شناختی سازمان‌دهی و بازیابی اطلاعات در سامانه اطلاعات علمی ایران (گنج)

یکی از چالش‌هایی ساماندهی و بازیابی اطلاعات ساختار زبانی متن است. ساختار زبانی متن شامل ویژگی‌های ساختارزی-نحوی-معنایی و لغوی متن است که تاثیر بسزایی بر دقت بازیابی دارد. در سامانه گنج به عنوان مثال، مسائل خط فارسی از جمله عواملی است که سبب کاهش دقت بازیابی می‌شود. در این پژوهش عوامل زبان‌شناختی سازمان‌دهی و بازیابی اطلاعات در سامانه اطلاعات علمی ایران (گنج) به روش میدانی شناسایی و بررسی شده و دسته‌بندی آنها بر اساس ادبیات موجود انجام می‌شود. سپس در مرحله بعد راهکارها یا راه‌حل‌هایی برای رفع یا بهبود آنها بر پایه مطالعات و روش‌های بکار گرفته شده در پژوهش‌های موجود در این زمینه ارائه می‌شود.

طراحی یک روش هوشمند تجزیه رشته‌های مرجع در زبان فارسی

یک رشته مرجع را می‌توان به عنوان مجموعه‌ای از مولفه‌ها مانند نام نویسندگان، عنوان، محل نشر، سال نشر، شماره صفحات و ... در نظر گرفت. در حالی که تجزیه رشته‌های مرجع موجود در انتهای یک مدرک علمی توسط کاربر انسانی به راحتی انجام‌پذیر است، تنوع موجود در فرمت‌های نگارش رشته‌های مرجع در کنار اشتباهات رخ داده توسط نویسندگان در نگارش این رشته‌ها، خودکارسازی انجام این عملیات را سخت کرده است. روش‌های زیادی برای خودکارسازی تجزیه رشته‌های مرجع ارایه شده است اما این روش‌ها وابسته به زبان بوده و امکان استفاده از یک روش ارایه شده برای یک زبان در زبانی دیگر منجر به نتایجی اشتباه می‌شود. تحقیقات صورت گرفته بیان‌گر این است که تاکنون هیچ روشی برای خودکارسازی تجزیه رشته‌های مرجع در زبان فارسی ارایه نشده است. با توجه به این مهم و نقش گسترده این مسأله در ساخت خودکار شبکه‌های استنادی مدارک علمی و فرایندهای بازیابی اطلاعات، در این طرح پژوهشی به این مسأله پرداخته شده است. در این راستا، پس از بررسی پیشینه مسئله، از روش یادگیری ماشین بردار پشتیبان به عنوان یک دسته‌بند چند دسته‌ای استفاده شده و یک روش هوشمند برای مسأله تجزیه رشته‌های مرجع در زبان فارسی ارایه شد. با توجه به اهمیت انتخاب ویژگی‌های مناسب برای استفاده در دسته‌بند ماشین بردار پشتیبان، در این پژوهش این مهم با توجه به ویژگی‌های استفاده شده در زبان انگلیسی و ویژگی‌های زبان فارسی و ارجاع‌دهی در این زبان انجام شد. روش پیشنهادی، پیاده‌سازی و با استفاده از مجموعه داده‌ای شامل ۹۲۲ رشته مرجع فارسی ایجاد شد و در نهایت آموزش داده شده و مورد آزمایش قرار گرفت. نتایج به دست آمده نشانگر مقدار ۹۲ درصد برای پارامترهای دقت، فراخوانی و معیار "ف" می‌باشد.

طراحی سامانه برچسب‌دهی به اجزای کلام برای متون فارسی

در این پژوهش به این مسئله پرداخته شد که آیا با رفع ابهام از برچسب نحوی هم‌نگاره‌های اسمی و صفتی مختوم به «-ی»، که فراوانی بالایی در پیکره‌های متنی فارسی دارند، کارایی یک سیستم برچسب‌زنی خودکار، افزایش می‌یابد و در نهایت می‌توان سامانه‌ای طراحی کرد که عمل برچسب‌دهی خودکار را با در نظر گرفتن رفع ابهام از برچسب هم‌نگاره‌های اسمی و صفتی مختوم به «-ی» در فارسی، با کارایی بهتری انجام دهد؟ سیستم مورد مطالعه در این پژوهش، سیستم «هضم» بود. نرم‌افزاری جهت رفع ابهام از برچسب نحوی هم‌نگاره‌های اسمی و صفتی مختوم به «-ی» در فارسی، تهیه شد که خود مبتنی بر الگوهای حساس به بافت نحوی است که بر اساس این الگوها می‌توان برچسب درست را به هم‌نگاره‌های مذکور اختصاص داد. ارزیابی کلی نرم‌افزار تهیه شده جهت رفع ابهام از برچسب نحوی

هم‌نگاره‌های اسمی و صفتی مختوم به «ی» در فارسی، نشان می‌دهد اگر تنها الگوهای حساس به بافت نحوی که تأثیر مثبت در برچسب‌زنی داشته‌اند را به برچسب‌زن «هضم» اضافه کنیم، صحت (Accuracy) کلی برچسب‌زن ۹۵,۶۹۱ درصد می‌شود که ۱,۳۴ درصد نسبت به حالتی که از تمام الگوهای حساس به بافت نحوی استفاده شود، بالاتر است.

ارایه روشی هوشمند برای استخراج کلیدواژه از مستندات علمی زبان فارسی بر اساس سیستم‌های پیشنهاددهنده

با توجه به افزایش حجم تولید و ثبت مستندات علمی، به خصوص در پایگاه گنج، نیاز است که فرایند نمایه‌سازی با سرعت بیشتری انجام شود. هم‌اکنون علاوه بر متوسط نرخ ثبت روزانه ۲۰۰ عنوان پایان‌نامه در روز در گنج، حدود ۱۴۰ هزار سند موجود است که هنوز نمایه نشده‌اند. در حال حاضر فرایند نمایه‌سازی در ایرانداک براساس دانش و تخصص نمایه‌سازان انجام می‌شود و روش‌های ماشینی برای پردازش خودکار متون و نمایه‌سازی آنها استفاده نمی‌شود. در بسیاری از پایگاه‌های علمی در دنیا از روش‌های ماشینی و خودکار در کلیه فعالیت‌های فرایند نمایه‌سازی یا بخشی از آنها استفاده می‌شود. تعدادی از این روش‌ها بر مبنای تحلیل آماری متون و استفاده از روش‌های یادگیری ماشین هستند، تعدادی بر مبنای تحلیل معنایی متون به واسطه اصطلاح‌نامه‌های تخصصی و هستان‌شناسی، و در تعدادی دیگر از این روش‌ها از تلفیق هر دو استفاده می‌شود. در این پژوهش، با هدف تسریع فرایند نمایه‌سازی، با استفاده از روش‌های یادگیری ماشین و تحلیل آماری متون روشی ارایه شد که از طریق آن برای نمایه‌سازی یک سند جدید، مجموعه‌ای از کلیدواژه‌های مرتبط به نمایه‌ساز پیشنهاد شود تا نمایه‌ساز سریع‌تر بتواند از بین آنها کلیدواژه‌های مناسب را انتخاب کند. پیشنهاد کلیدواژه به نمایه‌ساز، براساس رویکرد سیستم‌های پیشنهاددهنده و با اتکا به دانش ضمنی موجود در مجموعه مستندات نمایه شده قبلی صورت گرفت.

طراحی و ساخت سامانه‌ی استانداردساز و خطایاب متون فارسی

روزانه هزاران مستند متنی متنوع در حوزه‌های مختلف علمی بر روی وب جهان گستر قرار می‌گیرد. این مستندات می‌تواند شامل پایان‌نامه‌ها، مقاله‌ها، گزارش‌های علمی و مواردی از این قبیل باشد. نگارش متن این مستندات علمی جهت حفظ یکنواختی باید بر اساس اصول ثابت انجام گیرد، اما همواره به‌طور غیرعمدی دستخوش سلیقه‌های مختلفی در طول تاریخ می‌شود. اگرچه این تغییرات ناشی از پویا بودن زبان و خلاقیت ذهن بشری است، اما این پویایی و خلاقیت پردازش ماشینی متن را با چالش‌های متعددی روبرو می‌کند و دقت پردازش داده‌ها را به میزان چشمگیری پایین می‌آورد. علاوه بر تنوع نگارشی،

غلط‌های سهوی املائی نیز وجود دارد که فحوای گفتمانی متن را منحرف کرده و درک آن را با مشکل مواجه می‌کند. بنابراین، کلیه‌ی نویسه‌های متن باید به حالت استاندارد تبدیل شوند و عاری از هرگونه خطاهای املائی گردند. در این پژوهش، سامانه‌ای برای استانداردسازی و خطایابی متون علمی فارسی طراحی شده که متون نوشتاری علمی و تخصصی فارسی را به لحاظ صحت نگارشی و املائی بررسی کرده و متن را به شکل استاندارد درمی‌آورد.

عرضه قواعد سرواژه‌سازی در زبان فارسی

سرواژه‌سازی یکی از فرایندهای واژه‌سازی و یکی از انواع اختصارسازی است. در این نوع واژه‌سازی با حروف آغازین کلمات یا با حروف آغازین و برخی حروف دیگر کلمات و اصطلاحات صورت‌های جدیدی ساخته می‌شود. قواعد کلی سرواژه‌سازی در زبان فارسی شبیه به قواعد سرواژه‌سازی در زبان‌های غربی است اما برخی ویژگی‌های خط فارسی و ویژگی‌های آوایی زبان فارسی و حتی مسائل فرهنگی نیز قواعد سرواژه‌سازی در زبان فارسی موثراند. در این پژوهش تلاش شده با بررسی مسائل سرواژه‌سازی در زبان فارسی، قواعدی برای آن عرضه شود.

تعیین قواعد صوری همگون گروه‌های دستوری برای کاربرد در پردازش زبان

پردازش زبان یکی از حوزه‌های پژوهشی هوش مصنوعی است. ترجمه ماشینی و درک متن دو زمینه اصلی پژوهش در حوزه پردازش زبان‌اند. در دست داشتن قواعد صوری زبان و تبدیل آنها به قواعد و کدهای قراردادی از ملزومات و داده‌های اصلی در این گونه پژوهش‌ها است. استخراج قواعد دستوری گروه‌های دستوری و عرضه کدها و قواعد قراردادی مقدمه‌ی توصیف صوری دستور زبان است. تاکنون گروه‌های دستوری زبان فارسی به طور جداگانه بررسی و برای آنها کدها و قواعد قراردادی عرضه شده است. در این پژوهش کدهایی که تاکنون برای گروه‌های دستوری در طرح‌های جداگانه انتخاب شده بررسی و همگون شده و فهرست درهم کرد و جامع آنها عرضه شده است.

طراحی یک الگوریتم همانندجو برای تشخیص متون بازنویسی شده در زبان فارسی

همانندجویی ابزاری است که از آن برای تشخیص سرقت علمی/ادبی استفاده می‌شود. در یک روش همانندجویی، هدف تشخیص تمام قسمت‌های همانند موجود در یک متن مشکوک با توجه به تعدادی متن منبع احتمالی است. روش‌های زیادی برای همانندجویی ارائه شده اما از یک طرف، استفاده از روش‌های همانندجویی موجود برای سایر زبان‌ها به منظور همانندجویی در زبان فارسی مناسب نیست

و از طرف دیگر، اغلب روش‌های ارائه شده برای همانندجویی در زبان فارسی قادر به تشخیص متون بازنویسی شده نیستند. با توجه به این مهم، در این پژوهش، دو روش همانندجویی جدید با هدف تشخیص متون فارسی بازنویسی شده ارائه شده است. روش اول پیشنهادی روشی معنایی است و از لغت‌نامه جهت بررسی همانندی جملات متون استفاده می‌کند. روش دوم پیشنهادی روشی احتمالاتی است و از اطلاعات آماری به دست آمده از پیکره‌ای عظیم از متون برای همانندجویی استفاده می‌کند. علاوه بر این، در حالیکه در سایر روش‌های موجود، همانندی هر دو جمله از متون مورد نظر به صورت مستقل بررسی می‌شود، در روش‌های پیشنهادی همانندی جملات همسایه نیز در بررسی همانندی دو جمله در نظر گرفته شده است. نتایج پیاده‌سازی و آزمایشات صورت گرفته بر روی روش‌های پیشنهادی نشان می‌دهد که در حالیکه هر دو روش از کیفیت مناسب و تقریباً یکسانی برخوردار هستند، روش همانندجوی احتمالاتی پیشنهادی بسیار کاراتر است.

طراحی و پیاده‌سازی آزمایشی موتور سامانه هوشمند استخراج پژوهشگران مشابه

پژوهشگاه علوم و فناوری اطلاعات ایران با جمع‌آوری و پردازش تمام پایان‌نامه‌ها و رساله‌های وزارت علوم، تحقیقات و فناوری از جایگاه ممتاز و رو به رشدی برای دانشجویان، اساتید و مدیران برخوردار شده است. ارائه سرویس‌های ارزش افزوده برای سیاست‌گذاران و تصمیم‌گیرندگان حوزه علم و فناوری و امکان تجاری‌سازی دانش در جهت تثبیت نقش مرجعیت پژوهشگاه می‌تواند مفید باشد. استخراج تخصص پژوهشگران، مطالعه روند فعالیت علمی پژوهشگران، پژوهشگر در یک نگاه و ... نمونه‌ای از سرویس‌های ارزش افزوده با محوریت پژوهشگر می‌باشد. در پایگاه داده ایرانداک، یک پژوهشگر مشخص به اشکال مختلف در پایگاه داده ثبت شده است. دلایلی مانند چند املائی بودن برخی نام‌ها، عدم دقت دانشجو، غلط‌های املائی و ... را می‌توان برای علت گوناگونی نام پژوهشگر نام برد. از طرف دیگر تا زمانی که زیرساخت یکتاسازی پژوهشگران ترمیم نگردد طراحی و پیاده‌سازی سامانه‌های ارزش افزوده در حوزه پژوهشگران با شکست روبرو خواهد شد. با توجه به تعداد بالای پژوهشگران، در این پژوهش، ابتدا با استفاده از روش‌های زبان‌شناسی رایانشی و مفاهیم کلان داده یک معماری برای استخراج پژوهشگران مشابه طراحی شده و در ادامه نسخه آزمایشی آن پیاده‌سازی شده است.

بررسی ساختار گفتمانی چکیده‌های پایان‌نامه‌های فارسی در برخی رشته‌های علوم انسانی، فنی‌مهندسی و علوم پایه

رشته‌های مختلف علمی در نگارش چکیده دارای تنوع گفتمانی هستند. شناسایی این ساختار

گفتمانی برای سازماندهی اطلاعات در پایگاه‌های اطلاعاتی، تدوین استانداردهای چکیده‌نویسی و نیز سامانه‌های چکیده‌ساز خودکار متن از اهمیت برخوردار است. هدف از این پژوهش، تحلیل ساختار گفتمانی چکیده‌ی پایان‌نامه‌های فارسی در رشته‌های مختلف، تعیین حرکت‌های بلاغی آنها و به‌دست‌دادن تشابهات/تمایزات آنهاست. برای این منظور، سه شاخه‌ی علمی علوم انسانی، فنی‌مهندسی و علوم پایه تعیین و از هر شاخه نیز دو رشته‌ی نماینده برای تحلیل چکیده‌هایشان انتخاب شدند. رشته‌های نماینده عبارتند از: شیمی و فیزیک برای شاخه‌ی علوم پایه، مهندسی برق و مکانیک برای شاخه‌ی فنی‌مهندسی و فلسفه و زبان‌شناسی برای شاخه‌ی علوم انسانی. برای هر کدام از شش رشته، ۴۰ چکیده از پایگاه گنج ایرانداک جمع‌آوری و ساختار گفتمانی آنها در قالب حرکت‌های بلاغی و مدل پنج‌حرکتی هایلند کدگذاری شد. در مرحله‌ی بعد، حرکت‌های بلاغی هر کدام از رشته‌ها تحلیل و الگوی کلان برای شاخه‌های علمی به دست آمد. نتایج پژوهش نشان داده الگوی غالب برای اکثر رشته‌ها الگوی I-P-M-Pr و P-M-Pr است و در این میان دو رشته‌ی فلسفه و شیمی تفاوت‌هایی با دیگر رشته‌ها دارند.

شناسایی هویت نویسنده بر اساس زبان فردی: پژوهشی پیکره‌ای در زبان فارسی

زبان فردی به شیوه منحصر به فرد هر گویشور در به‌کارگیری عناصر زبانی از جمله واژه‌ها، ساختار دستوری و آواها گفته می‌شود که اخیراً موضوع پژوهش‌های رایانشی و زبان‌شناختی در شناسایی نویسنده متون فاقد هویت واقع شده است. یکی از مولفه‌های زبان‌شناختی که گفته می‌شود تجلی‌گاه زبان فردی واقع می‌شود، واژه‌های دستوری در زبان است. واژه‌های دستوری از آن جهت که به‌طور ناخودآگاه در تولید زبان به کار گرفته می‌شوند، مستقل از موضوع متن به کار می‌روند و بسامد بالایی در متون کوتاه دارند، همواره مورد توجه پژوهشگران در حوزه شناسایی سبک نویسنده بوده‌اند. در این پژوهش امکان تفکیک متون متعلق به یک نویسنده از متون دیگر با استفاده از واژه‌های دستوری زبان فارسی بررسی شده است. نتایج به دست آمده نشان می‌دهد که واژه‌های دستوری در زبان فارسی قابلیت تفکیک متون متعلق به یک نویسنده را دارند و عملکرد واژه‌های تک‌نگاشتی بهتر از دونگاشتی و سه‌نگاشتی‌ها در متون کم‌حجم است. همچنین نتایج پژوهش نشان می‌دهد که حجم کمینه متن برای شناسایی موفقیت‌آمیز نویسنده در متون فارسی حدود ۴۰۰۰ واژه بر اساس ۲۰ واژه دستوری پربسامد است.

رفع ابهام از برچسب نحوی هم‌نگاره‌های اسمی و صفتی فارسی

بررسی هم‌نگاره‌ها در پیکره‌های متنی فارسی نشان می‌دهد اکثر این هم‌نگاره‌ها، در اثر یکسان بودن نمود نوشتاری تکواژ یاء نکره، یاء اسم‌ساز، شناسهٔ دوم شخص مفرد، یاء صفت‌ساز و یاء متصل به گروه اسمی ایجاد شده‌اند. در صورت رفع ابهام از برجسب نحوی این هم‌نگاره‌ها (هم‌نگاره‌های اسمی و صفتی مختوم به «-ی»)، رفع ابهام معنایی از کلمات با سهولت بیشتری انجام خواهد شد. در این پژوهش ابتدا فهرست مبسوطی از هم‌نگاره‌های اسمی و صفتی مختوم به «-ی» تهیه شد؛ قواعد حساس به بافت نحوی جهت رفع ابهام از برجسب نحوی هم‌نگاره‌های مختوم به «-ی» استخراج شد. سپس جهت بررسی صحت قواعد مذکور، برنامهٔ ماشینی تهیه شد که صحت قواعد مسخرج از بررسی هم‌نگاره‌های مختوم به «-ی» را بررسی می‌کند. بررسی ماشینی صحت قواعد نشان می‌دهد، صحت بیش از نیمی از قواعد بالای ۷۰٪ است.

بررسی ساخت‌واژی هم‌نگاره‌های اسمی و صفتی به منظور کمک به برجسب‌دهی «اسم» به کلیدواژه‌ها در پیکره‌های علمی

در این تحقیق ساخت‌واژه اسم‌ها و صفت‌ها در فارسی به تفصیل بررسی شده است. نظام نوشتاری زبان فارسی نیز مورد بررسی دقیق قرار گرفت تا از این رهگذر بتوان به شناسایی انواع هم‌نگاره‌ها در زبان فارسی پرداخت. سپس، انواع هم‌نگاره‌ها در زبان فارسی مورد مطالعهٔ دقیق قرار گرفت و در نهایت از طریق جستجو به دو روش ماشینی و دستی، فهرست مبسوطی از هم‌نگاره‌ها از پیکره‌های «پیکرهٔ متنی زبان فارسی»، «پایگاه دادگان زبان فارسی» و «پیکرهٔ وابستگی نحوی زبان فارسی» تهیه شد. بررسی کلی هم‌نگاره‌ها در پیکره‌های مورد مطالعه نشان می‌دهد که بیشتر هم‌نگاره‌ها، فراوانی بالایی در پیکره‌های متنی فارسی دارند و اکثر آنها در اثر یکسان بودن نمود نوشتاری تکواژ یاء نکره، یاء اسم‌ساز، شناسهٔ دوم شخص مفرد، یاء صفت‌ساز و یاء متصل به گروه اسمی، ایجاد شده‌اند.

غنی‌سازی روابط معنایی هستان‌شناسی علوم پایه (شیمی)

هستان‌شناسی در واقع ابزار بازنمودن دانش در حوزه سازماندهی اطلاعات و هوش مصنوعی است و در دهه اخیر کاربرد آن در وب معنایی نیز مورد توجه قرار گرفته است. در تدوین هستان‌شناسی مفهومی، ایجاد یک ساختار سلسله‌مراتبی از مفاهیم، اولین ضرورت و ایجاد و تعیین نوع روابط بین مفاهیم در اولویت‌بعدی قرار دارد. در این پژوهش به منظور توسعه روابط معنایی هستان‌شناسی علوم پایه، که بر مبنای اصطلاح‌نامه‌های تخصصی تدوین شده در پژوهشگاه علوم و فناوری اطلاعات ایران تولید شده است، روشی نیمه‌خودکار برای غنی‌سازی روابط معنایی در حوزه شیمی ارائه شده است. در این روش ابتدا روابط سلسله‌مراتبی موجود در هستان‌شناسی مورد بازبینی و پالایش قرار می‌گیرد تا انواع

روابط سلسله‌مراتبی از یکدیگر متمایز شوند. در مرحله بعد مفاهیم موجود در هستان‌شناسی طبقه‌بندی و بر اساس روابط وابسته در اصطلاح‌نامه و الگوهای روابط معنایی استخراج شده از یک هستان‌شناسی سطح بالا، روابط معنایی بین مفاهیم تعیین شده و به هستان‌شناسی اضافه می‌گردد.

رویکردی پیکره‌ای و ساخت‌بنیاد به طبقه‌بندی معانی واژگانی در زبان فارسی

در این پژوهش با استفاده از رویکرد پیکره‌بنیاد، به بررسی افعال مرکب فارسی با جز فعلی «خوردن» پرداخته شده است. در فعل مرکب، جز فعلی نظیر «خوردن» که به آن فعل سبک نیز می‌گویند با عناصر غیرفعلی متعددی نظیر «ضربه»، «غذا»، «زمین» و غیره هم‌نشین می‌شود و این هم‌نشینی در افعال مرکب مختلف سبب می‌شود که فعل سبک «خوردن» معانی مختلفی در کنار اجزا غیرفعلی داشته باشد. آنچه مد نظر این پژوهش بوده، استفاده از روشی داده‌بنیاد برای طبقه‌بندی معانی مختلف فعل سبک «خوردن» بوده است. برای انجام پژوهش، در پیکره‌ای ۵۰ میلیون کلمه‌ای که به صورت تصادفی از پیکره همشهری نمونه‌گیری شده است مورد استفاده قرار گرفت. با استفاده از نرم‌افزار تحلیل پیکره‌ای AntConc عناصر غیرفعلی فعل سبک «شدن» استخراج شد. سپس ۳۵ فعل به ترتیب بسامد پیکره انتخاب و برای هر فعل ۲۰ نمونه کاربردی از پیکره استخراج و با استفاده از روش تحلیل نمای رفتاری، مشخصه‌های ساخت‌واژی، نحوی و معنایی آن برچسب‌دهی شد. سپس جدول نمای رفتاری افعال با استفاده از تحلیل خوشه‌ای سلسله‌مراتبی در نرم‌افزار R تجزیه و تحلیل شد. خوشه‌های معنایی به دست آمده از تحلیل خوشه‌ای نشان می‌دهد که معانی بر اساس بافت کاربردی افعال به خوشه‌های متمایز تقسیم می‌شوند و افعال مرکبی که اجزا غیرفعلی آنها به لحاظ حوزه معنایی مشترک هستند، در خوشه‌های مشابه و نزدیک به هم جای می‌گیرند.

تعیین قاعده وابسته‌های پسین گروه اسمی

منظور از پردازش زبان فراهم کردن تمهیداتی رایانه‌بنیاد است که براساس آن بتوان زبان را به شیوه‌ای ماشینی و در حد امکان بدون دخالت انسان فهم و درک کرد. در این پژوهش ابتدا پیشینه بررسی گروه اسمی و مشخصاً وابسته‌های پسین آن بررسی شده و سپس قواعد زبانی وابسته‌های پسین با ذکر مثال عرضه شده است. در نهایت این قواعد به صورت فرمول‌های قراردادی ارائه شده است. لازم به ذکر است که هدف از انجام این پژوهش استخراج قواعد صوری وابسته‌های پسین گروه اسمی زبان فارسی است که زیربنای پردازش گروه اسمی زبان فارسی خواهد بود.

ساخت پیکره گفتاری هجای CV زبان فارسی مبتنی بر الگوی نوایی گفتار برای استفاده در برنامه تبدیل متن به گفتار

در برنامه‌های تبدیل متن به گفتار معمولاً پردازش نوایی با تغییر فرکانس پایه و دیرش واحدهای بازسازی صورت می‌گیرد، که گفتار حاصل تا حدی ویژگی‌های نوایی جمله را دارد اما با توجه به آزمون‌های ادراکی و پژوهش‌های جدیدی که در زمینه بازسازی واحدهای گفتاری صورت گرفته به نظر می‌رسد بهتر است واحدهای بازسازی از بافت‌های نوایی مختلف و حساس به ساختار نوایی انتخاب شده و سپس برای تولد گفتار هم‌گذاری شوند. از این رو، در این پژوهش، تهیه واحدهای بازسازی یعنی هجا بر این اساس مدنظر قرار گرفت و پیکره نوشتاری هجای CV فارسی بر این اساس تهیه شده است. با توجه به پیچیدگی‌ها و دشواری‌های موجود، هجاهای مورد نظر تنها از چهاربافت نوایی (۱) هجای بی‌تکیه واژگانی؛ (۲) هجای تکیه بر واژگانی، تکیه زیر و بمی غیرهسته؛ (۳) هجای تکیه بر واژگانی، تکیه زیر و بمی هسته، نواخت مرزی خیزان استخراج شده و تولید هجاها در گفتار پیوسته و در درون جمله صورت گرفته است. ضبط و تهیه پیکره گفتاری هجاهای ساده و مرکب بر اساس سیگنال آکوستیکی و طیف نگاشت و با لحاظ معیارهای مشخص صورت گرفته و سپس هجاهای حاصل در پرونده‌های مختلف مبتنی بر جایگاه نوایی ذخیره شده‌اند. داده‌های حاصل می‌توانند برای استفاده در برنامه تبدیل متن به گفتار مورد استفاده قرار گیرند.

طراحی هستی‌شناسی علوم پایه و پایگاه اطلاعات مفهوم- بنیان آن

تعامل میان هستی‌شناسی و زبان‌های طبیعی مقدم بر فرایند پردازش زبان طبیعی و مهندسی دانش است. هستی‌شناسی‌ها که عمدتاً از الگوهای معناشناختی بهره می‌برند وجوه مشترکی با واژگان دارند. یک مجموعه واژگان اصطلاحات را براساس مفاهیم مرتبط سازماندهی می‌کند، در حالی که هستی‌شناسی‌ها مفهوم‌سازی کرده و روابط مفاهیم را تعیین می‌نماید. یک واژگان مشترک پیش‌زمینه و ابزاری است در به اشتراک گذاشتن دانش جمعی و یک هستی‌شناسی مشترک علاوه بر این ابزاری است در به اشتراک گذاشتن دانش از طریق فناوری اطلاعات. در ساخت الگوهای مفهومی و ساختار زبانی، این حوزه زبان‌شناسی رایانه‌ای است که در ترسیم روابط میان مفاهیم و واژگان به منظور کاربرد در مهندسی دانش و سازماندهی اطلاعات نقش ایفا می‌کند. به منظور شکل‌گیری هستی‌شناسی مشترک به تلفیق و همسان‌سازی اصطلاح‌نامه‌های موجود در پژوهشگاه علوم و فناوری اطلاعات پرداخته، مغایرت‌ها و هم‌پوشانی‌های واژگان علوم پایه را برطرف و شبکه مفاهیم این گرایش را با

زیرمجموعه‌های آن به شکل هستی‌شناسی واژگانی علوم پایه طراحی نمودیم تا پس از تبدیل روابط اصطلاح‌نامه‌ای به مفهوم-بنیاد به زبان OWL2 و بر اساس استاندارد ۲۵۹۶۴ در پایگاه اطلاعات مفهوم-بنیاد به عنوان ابزاری قوی جهت ذخیره و بازیابی اطلاعات به کار گرفته شود.

شناسایی طرح‌های مرتبط با تحقق مترجم ماشینی در زبان فارسی

ترجمه ماشینی زیر شاخه‌ای از زبان‌شناسی محاسباتی است که عبارت است از ترجمه متنی از یک زبان طبیعی به زبانی دیگر توسط کامپیوتر. ترجمه ماشینی از دیرباز دغدغه بشر بوده و به مرور با گذشت سالیان متمادی توانسته به نقطه‌ای برسد که امروز شاهد آن هستیم. در این پژوهش، در ابتدا به بررسی پیشینه مساله پرداخته شده و پژوهش‌های انجام شده توسط دیگر متخصصان این حوزه و برخی نرم‌افزارهای موجود معرفی شده است. در ادامه، به هدف اصلی این پژوهش پرداخته و طرح کلی ساخت نرم‌افزار مترجم فارسی ارائه شده است.

بهبود سیستم بازشناسی اعداد تلفنی با استفاده از پارامترهای آکوستیکی آواهای زبان فارسی

در گفتار تلفنی عوامل متعددی سبب از بین رفتن تمایزات واجی و در نتیجه ابهام‌آوایی می‌شوند. اما این تمایزات توسط برخی از پارامترهای آکوستیکی باقیمانده در سیگنال گفتار، انتقال پیدا می‌کنند. تحلیل ۱۸۱ نمونه آوایی که توسط ۱۳ زن و ۱۴ مرد در بافت‌های مختلف آوایی تولید شده‌اند، نشان می‌دهد که پارامترهای نوفه سایش شامل دیرش سایش و مرکز ثقل طیف و نیز اطلاعات آوایی واکه شامل دیرش واکه، بسامد پایانه F2 تمایز واجی [se]-[sel] را حتی در نمونه‌های کاهش‌یافته‌ای مانند [seF] در سیگنال گفتار انتقال می‌دهند. از آنجا که پارامترهای واکه، سرنخ‌های دقیق‌تر و بهتری برای انتقال تمایز واجی [se]-[sel] در سیگنال تلفنی هستند، سه پارامتر دیرش واکه، بسامد پایانه F2 و شیب‌گذر F2 در ۸۱۱ نمونه آوایی (۴ زن و ۹ مرد) مورد اندازه‌گیری قرار گرفت و به عنوان داده‌های تست برای افزایش دقت سیستم بازشناسی استفاده شد. برای افزایش دقت بازشناسی از روش تلفیق دو شیوه بازشناسی بر اساس مدل مخفی مارکوف و ضرایب مل‌کپستروم و نیز بازشناسی بر اساس پارامترهای استخراج شده از واکه استفاده شد. برای ترکیب نتایج حاصل از دو روش فوق، از معیار اطمینان استفاده شد. یکی از روش‌هایی که می‌توان برای سیستم مبتنی بر پارامتر، معیار اطمینان تعریف کرد، استفاده از درجه شباهت نرمال‌شده‌ای است که با استفاده از توزیع گوسی به دست می‌آید. دقت حدود یک درصد بیشتر از سیستم بازشناسی مبتنی بر HMM ۳۳/۹۳ درصد است و به ۳۱/۹۴ درصد می‌رسد. بنابراین می‌توان گفت روش ترکیبی مبتنی بر معیار اطمینان خطای بازشناسی ارقام سه و صفر

به طور نسبی حدود ۱۵ درصد کاهش داده است.

طراحی مدل مفهومی یکپارچه نمایه‌سازی ماشینی منابع فارسی

در این پژوهش، در ابتدا به بررسی رابطه بین نمایه‌سازی و بازیابی اطلاعاتی (جست‌وجو)، مباحث نظری و سیر تحول نمایه‌سازی ماشینی از دیدگاه تئوریک، نمایه‌سازی ماشینی به صورت عملیاتی و از نگاه فرایندی، پروژه‌های موفق انجام شده در دنیا و برای زبان‌های مختلف و نیازمندی‌های نمایه‌ساز ماشینی در زبان فارسی پرداخته شده است. در ادامه، اجزای یک نمایه‌ساز زبان فارسی تشریح گردیده و مدل مفهومی نمایه‌ساز ماشینی به صورت دقیق ترسیم شده و روابط بین اجزاء، جزئیات زیر سیستم‌ها و مدل منطقی اولیه سیستم طراحی شده است. در انتها نیازمندی‌های توسعه نمایه‌ساز ماشینی برای پژوهشگاه علوم و فناوری اطلاعات ایران مورد بحث واقع شده و وضعیت پروژه‌هایی که در حال حاضر انجام شده و یا تجربیاتی که وجود داشته و می‌تواند در اجرای این پروژه مفید باشد مورد توجه قرار گرفته است.

تعیین قاعده‌های صوری وابسته‌های پیشین گروه اسمی و کاربرد آن در ترجمه ماشینی

پردازش زبان یکی از حوزه‌های پژوهشی هوش مصنوعی بوده و ترجمه ماشینی نوعی پردازش زبان از زبان مبدا به زبان مقصد است. یکی از ملزومات ترجمه ماشینی استخراج قواعد صوری دستور زبان است که در این راستا باید برای تمامی حوزه‌های نحو زبان قواعدی صوری استخراج کرد. یکی از حوزه‌های نحو مربوط به گروه اسمی است. گروه اسمی مهمترین سازه در پردازش دستور زبان فارسی است. گروه اسمی شامل قواعد وابسته‌های پیشین و قواعد وابسته‌های پسین است. در این پژوهش به تعیین قاعده‌های صوری وابسته‌های پیشین گروه اسمی پرداخته و کاربرد آن در ترجمه ماشینی مورد بررسی قرار گرفته است.

فعل‌شکن فارسی

قواعد صوری دستور زبان یکی از ملزومات پردازش زبان است. پردازش نحوی دستور زبان یکی از حوزه‌های پردازش زبان است. برای پردازش نحوی باید قواعد صوری حوزه‌های مختلف دستور زبان را استخراج کرد. یکی از حوزه‌های نحو، حوزه افعال است. در این پژوهش به استخراج قواعد برنامه‌ای رایانه‌ای که قادر به تشخیص و تجزیه فعل‌های ساده بوده پرداخته شده است. برنامه خروجی که فعل‌شکن فارسی نامیده شده برنامه‌ای رایانه‌ای است که صیغگان فعل‌های زبان فارسی را تجزیه و

شناسایی می‌کند و قادر به شناسایی و تجزیه ساختمان انواع فعل‌های ساده با قاعده و بی‌قاعده زبان فارسی است. در این طرح فعل‌های مرکب و پیشوندی بررسی نشده‌اند. زیرا هدف همکرد در فعل‌های مرکب و عنصر غیرفعلی در فعل‌های پیشوندی، آن‌چه باقی می‌ماند در واقع یک فعل ساده است.

عددشکن فارسی

زبان فارسی همچون بسیاری دیگر از زبان‌ها دارای قواعد ترکیب اعداد است. قواعد ترکیب اعداد بخشی از گروه اسمی در زبان فارسی‌اند. گروه اسمی دارای وابسته‌های پیشین و وابسته‌های پسینی است که قبل و بعد از هسته گروه اسمی قرار می‌گیرند. اعداد و صور ترکیبی آنها هم قبل از هسته گروه اسمی و هم بعد از آن می‌آیند. عددشکن فارسی برنامه‌ای رایانه‌ای است که صورت مکتوب اعداد و ترکیبات آن در زبان فارسی را به شکل عددی آن نمایش می‌دهد. این نرم‌افزار حافظه‌ای دارد که قواعد ترکیب اعداد فارسی در آن تعبیه شده است. نرم‌افزار شکل مکتوب را می‌خواند و آن را با قواعدی که در حافظه دارد منطبق می‌کند و سپس آن را به صورت اعداد عرضه می‌کند.

واژه‌شکن فارسی

واژه‌شکن فارسی برنامه‌ای رایانه‌ای است که کلمات پیچیده را به اجزاء سازنده‌اش تجزیه و مقوله دستوری و معنایی آن را مشخص می‌کند. این نرم‌افزار حافظه‌ای دارد که کلمات بسیط و تکواژها و قواعد ساخت کلمات در آن قرار دارد. نرم‌افزار ابتداء کلمه‌ای را که به آن عرضه شده در حافظه جستجو و مقوله دستوری و معنایی آن را می‌یابد. چنانچه صورت کامل کلمه موردنظر در حافظه وجود نداشته باشد نرم‌افزار در صدد تجزیه آن برآمده و در صورتی که اجزاء کلمه مفروض قاعده‌مند باشد، مقوله آن را مشخص می‌کند. چنانچه کلمه غیر دستوری باشد، نرم‌افزار تنها به تجزیه کلمه اکتفاء می‌کند.

طراحی الگوی نمایه استنادی نشریات فارسی ایران

نمایه استنادی مهمترین و اصلی‌ترین ابزار در تحلیل استنادی می‌باشد، از این طریق می‌توان ضمن آنکه امکان بازیاب اطلاعات را فراهم آورد، به ارزیابی کمی و کیفی تولیدات علمی، مطالعه تاریخ و ساختار علم، شناخت رابطه بین استنادها و ترسیم ساختار موضوعی رشته‌های علمی دست یافت. این پژوهش به مطالعه و طراحی یک الگو جهت ایجاد نمایه استنادی نشریات فارسی ایران پرداخته است. نمایه استنادی یکی از بسترهایی است که می‌تواند به عرضه کتاب‌شناختی مقالات کلیدی و برجسته علمی دست بزند و علاوه بر آن ارزیابی پژوهش و پژوهشگران را در میان انبوه اطلاعات امکان‌پذیر سازد. در این پژوهش به مطالعه و بررسی تعاریف و مفاهیم نمایه استنادی، شاخص‌های انتخاب نشریات

و سوابق سایر کشورها و ایران پرداخته شده است. علاوه بر آن ساختار پایگاه‌های استنادی موسسه اطلاعات علمی آمریکا مورد بررسی قرار گرفته و فرایند تولید و انجام کار در این نمایه‌های استناد ارائه گردیده است. نهایتاً بر اساس مجموعه اطلاعات جمع‌آوری شده الگوی ایجاد یک نمایه استنادی برای نشریات فارسی در ایران ارائه شده است.

فرایند نمایه‌سازی

استفاده موثر از اطلاعات به سازماندهی آن نیاز دارد و سازماندهی به فرایندهای متنوع بازنمایی محتوای مدرک برای بهبود جست‌وجو و بازیابی اطلاعات همراه است. یکی از مهمترین شیوه‌های بازنمایی محتوا، نمایه‌سازی است که با تخصیص توصیفگرها به مدارک، محتوای موضوعی آنها را نشان می‌دهد. نمایه‌سازی، یک مدرک را در میان انبوهی از منابع موجود در پایگاه‌ها و مخازن اطلاعات قابل بازیابی می‌سازد. نمایه‌سازی فرایندی پیچیده است و ارائه تعاریف، اصول، روش‌ها، ماهیت انواع نمایه‌ها، و روش‌های عملی نمایه‌سازی می‌تواند به استفاده بهینه از اطلاعات موجود کمک شایان نماید. این پژوهش ابتدا به تعریف موجزی از اطلاعات، و سازماندهی اطلاعات و روش‌های آن پرداخته است. سپس بر نمایه‌سازی به عنوان یکی از روش‌های سازماندهی تاکید و تعاریف نمایه و نمایه‌سازی، انواع آن، ساختار نمایه، و مسائل مربوط به سیاست‌گذاری نمایه‌سازی را ارائه کرده است. در ادامه درباره نحوه ایجاد نمایه، ویژگی‌ها و فرایند نمایه‌سازی، و همچنین انواع زبان‌های نمایه‌سازی بحث شده است. در پایان ضمن ارائه مباحثی در ارزیابی نمایه‌سازی، نمایه‌سازی خودکار، و نقش نمایه در بازیابی اطلاعات پیشنهادهایی در زمینه بهبود نمایه‌ها مطرح شده است.

بهینه‌سازی اطلاعات علمی شیمی و مهندسی شیمی (سمینارها) در بانک جامع اطلاعات پژوهشگاه اطلاعات و مدارک علمی ایران و طراحی نرم‌افزار بانک ترکیبات شیمیایی

در این پژوهش اطلاعات علمی ۲۷۱۴ رکورد اطلاعاتی مربوط به سمینارهای شیمی و مهندسی شیمی که در فاصله زمانی سال‌های ۱۳۶۸ تا ۱۳۷۹ هجری شمسی در کشور برگزار شده‌اند به صورت یک مجموعه واحد و دارای نمایه‌نامه‌های تخصصی به صورت چاپی و نیز نرم‌افزاری ارائه شده است. کلیه اطلاعات مورد بازبینی، ویرایش و بازپردازش تخصصی قرار گرفتند. بخصوص کلیدواژه‌ها مورد استاندارد کردن، یکسان‌سازی، ارجاع دهی و معادل‌یابی قرار گرفتند و در نهایت در یک مجلد به چاپ رسیدند، تا مورد استفاده دانش‌پژوهان و محققان در بخش صنعت و دانشگاه جهت تعیین اولویت‌های تحقیقاتی و جلوگیری از دوباره‌کاری قرار گیرد. کنترل نامگذاری‌ها با توجه به استانداردهای ایوپاک انجام شده است. همچنین شیوه بازیابی اطلاعات شیمی افزون بر جستجوی موضوعی با استفاده از

فرمول‌ها، اسامی ترکیبات شیمیایی و نیز شماره‌های ثبت شیمیایی عملی شده است. به منظور تقویت و اصلاح در ساختار تشکیلات نمایه‌سازی و چکیده‌نویسی در پژوهشگاه، دو نفر کارشناس شیمی مورد آموزش قرار گرفتند. طراحی نرم‌افزار تخصصی شیمی در ثبت، ذخیره و بازیابی اطلاعات ترکیبات شیمیایی، ایجاد نمایه‌نامه‌های تخصصی شیمی، طراحی کاربرگ‌های استاندارد ورود اطلاعات شیمی و استاندارد کردن شیوه‌های ذخیره و بازیابی اطلاعات شیمی از موارد مورد توجه در این پژوهش است. همچنین فهرست‌های تخصصی و واژه‌های یکدست شده در این طرح می‌توانند توسط دانشجویان، اساتید، محققان، نمایه‌سازان، مترجمان و فرهنگ‌نویسان مورد استفاده قرار گیرد.

بررسی اطلاعات علمی شیمی و مهندسی شیمی در بانک جامع اطلاعات مرکز به منظور بهینه‌سازی آنها

در این پژوهش اطلاعات علمی ۱۳۰۰ پایان‌نامه مربوط به رشته‌های شیمی و مهندسی شیمی در مقطع کارشناسی ارشد و دکتری از سال ۱۳۴۱ تا سال ۱۳۷۵، مورد بازبینی، ویرایش و باز پردازش تخصصی قرار گرفته است. به خصوص کلیدواژه‌ها، مورد استاندارد کردن، یکسان‌سازی، ارجاع‌دهی و معادل‌یابی قرار گرفته و در نهایت در یک مجلد به چاپ رسیده، تا محققانی که در بخش پژوهش و صنعت شیمی فعال هستند جهت تعیین اولویت‌های تحقیقاتی و جلوگیری از دوباره کاری آن را مورد استفاده قرار دهند. کنترل نامگذاری‌ها با توجه به استانداردهای ایوپاک انجام شده است. همچنین شیوه بازیابی اطلاعات شیمی با توجه به فرمول‌ها و اسامی ترکیبات شیمیایی عملی شده است. ضرورت تغییر در ساختار تشکیلات نمایه‌سازی و چکیده‌نویسی در مرکز، نمایه‌سازی و چکیده‌نویسی مدارک شیمی توسط متخصصان موضوعی، استفاده از نرم‌افزارهای تخصصی شیمی در ثبت و ذخیره و بازیابی اطلاعات شیمی، ایجاد نمایه‌نامه‌های تخصصی شیمی، استاندارد کردن شیوه‌های ذخیره و بازیابی اطلاعات شیمی و نهایتاً ایجاد پایگاه تخصصی شیمی از مهمترین مطالب مورد توجه در این پژوهش است. همچنین فهرست‌های تخصصی و واژه‌های یکدست شده در این پژوهش می‌توانند بطور وسیعی توسط دانشجویان، اساتید، محققان، نمایه‌سازان، مترجمان و فرهنگ‌نویسان مورد استفاده قرار گیرد.

یکسان‌سازی شیوه رسم‌الخط اسامی ترکیبات شیمیایی در زبان فارسی

ضبط چندگانه صورت‌هایی واحد با خط فارسی باعث ایجاد مشکلاتی در کار تعلیم و تربیت می‌شود و وقت بسیار برای یکسان کردن شیوه نگارش متون از دست می‌رود. یک‌دست کردن خط از ابزار ضروری پردازش زبان نیز هست. یکی دیگر از پیامدهای همگون نبودن رسم‌الخط، ایجاد نابسامانی در کار اطلاع‌رسانی علوم است. ضبط اعلام و کلیدواژه‌ها به شکل‌های مختلف از مشکلات رایج در

کتابخانه‌ها و مراکز اطلاع‌رسانی است. شیمی یکی از علومی است که اصطلاحات آن به گونه‌هایی متفاوت ثبت می‌شود. و مشکلات سابق‌الذکر بیشتر در آن نمایان است. در این پژوهش سعی شده است که شیوه نگارش اسامی ترکیبات شیمیایی و بالاخص ترکیبات آلی آن در زبان فارسی و معضلات مربوط به آن بررسی و الگوهای کلی برای یکسان‌نویسی آنها پیشنهاد شود. این پیشنهادات منطبق با اصول و ضوابط رسم‌الخط پیشنهادی فرهنگستان زبان و ادب فارسی است.

امکان‌سنجی ساخت پیکره از داده‌های علمی پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)

پیکره‌های زبانی و واکاوی پیکره بنیاد زبان در زبان‌شناسی پیکره‌ای نقشی کلیدی دارند. ساخت پیکره‌ها کاری پیچیده، زمان‌بر، دارای گام‌های گوناگون، و میان‌رشته است. بنابراین برای کارایی و اثربخشی در ساخت آن‌ها نیاز است که پیش از آغاز کار، امکان‌سنجی انجام شود. امکان‌سنجی؛ گام‌ها، هزینه‌ها، نیروی انسانی، حقوق مادی و معنوی، و مانند آن‌ها را برای یک پروژه، بررسی و به مدیریت پروژه کمک می‌کند تا آن را با آمادگی و آینده‌نگری بیشتری به سرانجام برساند. از آنجایی که مدلی برای امکان‌سنجی ساخت پیکره‌ی متنی همچون گونه‌ای از پیکره‌ها در پژوهش‌های پیشین نبود، این پژوهش در پی ساخت چنین مدلی انجام شد. برای این کار، پس از بررسی پژوهش‌های پیشین، شاخص‌های کلی ساخت پیکره‌های زبانی استخراج شد. سپس، فرآیند ساخت پیکره به تفصیل مورد بررسی قرار گرفت. پس از بررسی پژوهش‌های پیشین در زمینه امکان‌سنجی، بر پایه اطلاعات به دست آمده از مطالعات امکان‌سنجی، بررسی شاخص‌های ساخت پیکره و فرآیند ساخت پیکره، ابعاد و مؤلفه‌های امکان‌سنجی ساخت پیکره استخراج شد و مدلی کلی برای ساخت پیکره پیشنهاد شد. سپس با روش دلفی و در دو دور اعتباریابی، مدل پایانی به دست آمد. این مدل دارای هفت بُعد فنی، اقتصادی، زمان‌بندی، قانونی-حقوقی، عملیاتی، تأمین مالی، و بازاریابی و روی هم رفته ۳۳ مؤلفه است. این مدل و کاربرد آن در پروژه‌های ساخت پیکره‌ی متنی می‌تواند کمک بزرگی برای پژوهشگران و سازمان‌ها باشد. سپس در آخرین مرحله از پژوهش، مدل امکان‌سنجی ساخت پیکره در پنلی متشکل از متخصصان موضوعی و اجرایی در ایرانداک به شور گذاشته و اعتباریابی شد و مدل پایانی برای ساخت پیکره از داده‌های علمی ایرانداک به دست آمد و در پایان، امکان‌سنجی ساخت پیکره از داده‌های ایرانداک نیز انجام شد و چگونگی پرداختن به ابعاد و مؤلفه‌های مدل مشخص شد.

طراحی روشی هوشمند برای دسته‌بندی نوشتارهای علمی فارسی

با توجه به حجم بالای متون علمی و سرعت پایین کنترل اقلام اطلاعاتی در سامانه ثبت و در نتیجه عدم نمایش به روز در گنج، افزایش سرعت کنترل اقلام اطلاعاتی از موارد اولویت دار سامانه ثبت است. به عنوان مثال حدود ۱۴۰ هزار متون علمی همچنان نمایه نشده و با توجه به سیل پایان نامه های جدید، استفاده از روش هایی که می تواند کمکی به افزایش سرعت نمایه سازی کنند از اولویت های سامانه ثبت می باشد. پایان نامه و رساله هایی که در پایگاه داده ثبت پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک) ثبت می گردد دارای فراداده موضوع اصلی مانند علوم انسانی، فنی و مهندسی، علوم پایه، هنر و معماری، علوم پزشکی و غیره می باشند. استخراج هوشمند حوزه پارساها بوسیله الگوریتم های یادگیری ماشین باعث افزایش سرعت فرایند نمایه سازی و افزایش کیفیت داده های پایگاه داده گنج شود. در این طرح پژوهشی هدف دسته بندی هوشمند محتوایی متون به یکی از ۵ دسته موضوع اصلی علوم انسانی، فنی و مهندسی، علوم پایه، هنر و معماری، علوم پزشکی بوده است. به عبارت دیگر افزایش سرعت نمایه سازی بوسیله خود کار سازی تایید یا عدم تایید موضوع اصلی ثبت شده از طرف پژوهشگر. در این پژوهش ابتدا پیش پردازش معمول در پردازش زبان طبیعی انجام شده و ویژگی ها متمایز کننده استخراج شده است. در ادامه با استفاده از دسته بندی های متنوع، محتوای پایان نامه ها به پنج دسته آموزش داده شد. نتایج ارزیابی ها بر روی متون پایان نامه و رساله های فارسی نشان می دهد این روش با ارزیابی با متون بیشتر می تواند با دقت و سرعت خیلی خوبی موضوع اصلی پارساها را مشخص کند.

بررسی موتور جستجوی گنج به منظور بهبود بازیابی اطلاعات در زبان فارسی

بازیابی اطلاعات به معنای یافتن محتوایی با طبیعت غیر ساخت یافته (معمولاً متن) از مجموعه هایی به منظور برآورده ساختن نیازهای اطلاعاتی کاربران است. با بهره گیری از روش های بازیابی اطلاعات موجود در ادبیات پلتفرم های متن باز متعددی طراحی شده اند. کتابخانه آپاچی لوسن و موتور جستجوی آپاچی سولر از این نوع پلتفرم ها هستند. سامانه گنج نیز به منظور جست و جو در اسناد و مدارک ذخیره شده در پایگاه داده خود، از این کتابخانه و موتور جست و جو به عنوان هسته اصلی فرآیند بازیابی اطلاعات استفاده می کند. علی رغم کاربردهای گسترده، سولر دارای کاستی هایی بوده، تاکنون نیز هیچ گونه پژوهشی در زمینه شناسایی و واکاوی این کتابخانه و موتور جست و جو در پژوهشگاه انجام نشده و تنها عملکرد رابط کاربری آن و مشکلات موجود بررسی شده است. لذا علی رغم ضرورت بررسی کارایی رابط کاربری گنج، تازمانی که هسته اصلی آن (خصوصاً با تمرکز بر زبان فارسی) بررسی و اصلاح نشود نمی توان انتظار داشت عملکرد سامانه مطلوب باشد. لذا ضروری است تا با بررسی محدودیت ها و مشکلات موجود در روش های مورد استفاده در موتور جست و جو آپاچی سولر (از

منظر بازیابی اطلاعات) دیدگاهی کلان نسبت به آن در پژوهشگاه ایجاد شده تا از این طریق بتوان پیشنهادهایی به منظور ارتقا عملکرد هسته اصلی فرآیند بازیابی اطلاعات سامانه گنج ارائه داد. بر این اساس، در این طرح دانش فنی درمورد این موتور جست‌وجو و روش‌های مورد استفاده آن کسب شده و راهکارهایی به منظور ارتقای نتایج بازیابی در زبان فارسی و در سامانه گنج ارائه شد. راهکارهای پیشنهاد شده در قالب دو دسته راهکارهای نیازمند پژوهش و راهکارهای فنی و عملیاتی دسته‌بندی شدند.

کتاب

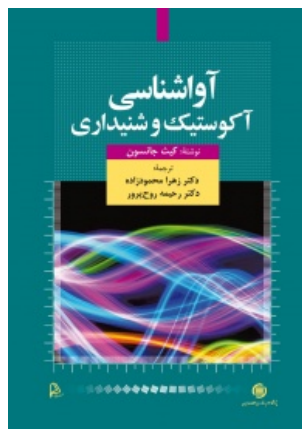
تاریخچه آواشناسی و سهم ایرانیان

آواشناسی عمدتاً در چارچوب مطالعه زبان بررسی شده است. به این سبب دستاوردهای آواشناسی در این کتاب از میان آثاری که درباره مسائل کلی زبان است، استخراج و با تسلسلی تاریخی عرضه شده است. محدوده زمانی آثار از حدود شش قرن قبل از میلاد تا قرن پانزدهم میلادی است. نویسنده سعی کرده داور عادل باشد و سهم واقعی دانشمندان ایرانی را در آواشناسی مشخص کند، نه آنکه دانشمندان هموطنش را از سر احساس در بهترین جای یک قصه علمی تخیلی قرار دهد. در این کتاب آثار و بقایای آواشناسی از خرابه‌های هند و یونان و روم و سرزمین‌های اسلامی جمع شده و در بافت اجتماعی دیاری که دانشمندان در آن به بررسی این علم پرداخته‌اند، عرضه شده است.



آواشناسی آکوستیک و شنیداری

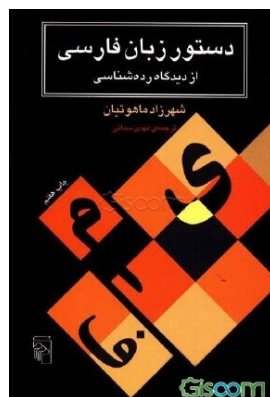
این کتاب مقدمه‌ای کوتاه و غیرفنی (مناسب به عنوان یک منبع مکمل برای درس آواشناسی عمومی یا علم گفتار) در چهار مبحث مهم آواشناسی آکوستیک شامل: ۱) ویژگی‌های آکوستیکی طبقات عمده آواهای گفتاری؛ ۲) نظریه آکوستیکی تولید گفتار؛ ۳) بازنمایی شنیداری گفتار؛ و ۴) این کتاب برای دانشجویان در درس‌های مقدماتی آواشناسی زبان‌شناسی، علوم گفتاری و شنیداری و شاخه‌هایی از مهندسی برق و روان‌شناسی شناختی مورد استفاده خواهد بود. پنج فصل اول به معرفی مبانی آکوستیک، نظریه آکوستیکی تولید گفتار، پردازش دیجیتالی سیگنال، شنوایی و درک گفتار می‌پردازد و در چهار فصل باقیمانده، طبقات عمده آوایی مورد مطالعه قرار می‌گیرد و مشخصه‌های آکوستیکی آواهای گفتاری طبق پیش‌بینی نظریه آکوستیکی تولید گفتار، ویژگی‌های شنیداری و مشخصه‌های درکی آنها مرور می‌شود. در پایان هر فصل، فهرستی از منابع پیشنهادی برای مطالعه بیشتر به همراه تعدادی تمرین آمده است. تمرین‌ها با تأکید بر اصطلاحاتی که در هر فصل با حروف سیاه مشخص گردیده (و در قسمت «واژگان تخصصی» فهرست شده‌اند)، خواننده را به به کار بردن مفاهیم مطرح شده در فصل ترغیب می‌نمایند.



دستور زبان فارسی از دیدگاه رده‌شناسی

این کتاب نخستین دستور زبان فارسی است که از دیدگاه رده‌شناسی زبان که یکی از گرایش‌های عمده در زبان‌شناسی امروز است، نوشته شده است. اصل کتاب جزو مجموعه‌ای است که به قصد فراهم آوردن داده‌هایی واحد و دارای نظام برای محققان زبان‌شناسی نظری و پژوهشگران تدوین شده و به این هدف تاکنون شصت زبان بررسی و برای آنها دستور نوشته شده است. این دستور شامل سه بخش صرف، نحو و واج‌شناسی است و در مباحث آن مقایسه‌هایی نیز بین ویژگی‌های دستوری فارسی و دیگر زبان‌ها انجام شده است. کتاب، دستوری توصیفی است و در چارچوب بررسی «هم‌زمانی»، زبان

محاوره‌ای فارسی تهرانی افراد تحصیل کرده را انتخاب و آن را توصیف کرده است. نکته‌های بدیع، ساختار، هدف خاص و مباحث جدید کتاب بر جامعیت و سودمندی آن افزوده است.



ساختواژه، اصطلاح‌شناسی و مهندسی دانش

مهندسی دانش، به منظور سازماندهی اطلاعات علمی، به برنامه‌ریزی رایانه‌ای و ایجاد شبکه‌ی اطلاعاتی می‌پردازد. ساماندهی نام‌گذاری مفاهیم علمی، یکی از مهم‌ترین مباحث علمی اصطلاح‌شناسی محسوب می‌گردد که در آن، «مفهوم» به عنوان واحد دانش فنی، تعریف شده است. به منظور ترسیم روابط مفهومی یک رشته تخصصی باید نظام مفاهیم زیررشته‌های مرتبط با آن را، در یک ساختار کلی تلفیق نمود. بر این اساس، اصطلاح‌نامه‌های تخصصی تدوین می‌شوند و در نهایت با یکدیگر ادغام می‌گردند. فرااصطلاح‌نامه نوعی اصطلاح‌نامه کلان است که موجب تلفیق اصطلاح‌نامه‌های متعدد موجود در سطح ملی و بین‌المللی می‌گردد و در بازایی، ارجاعات متقابل میان پایگاه‌های اطلاعات علمی گوناگون برقرار می‌نماید. در اثر حاضر با استفاده از بانک اطلاعات واژگان (تلفیق واژگان علمی اصطلاح‌نامه‌های علوم پایه) پژوهشگاه علوم و فناوری اطلاعات ایران که به عنوان پیکره زبانی انتخاب گردیده است و با توجه به نظریه ساختوازی و معناشناسی واژگانی (لیبر ۲۰۰۴) این واژگان مورد تجزیه و تحلیل آماری قرار گرفته‌اند و علاوه بر تعیین فراوانی هر یک از عناصر واژگان مرکب، مشتق و مرکب مشتق، الگوهای ساختوازی اصطلاحات علمی استخراج شده است. سپس، با تعیین روابط معنایی آنها، زمینه طراحی هستی‌شناسی و شبکه‌ی واژگانی این حوزه علمی برای برنامه‌ریزی آموزشی و پژوهشی و در نهایت، طراحی نظام مفاهیم شبکه‌ی علمی کشور فراهم گردیده است.

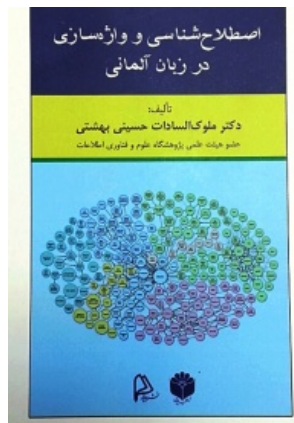


سامانه استانداردسازی و خطایاب متون علمی فارسی

روزانه هزاران مستند متنی متنوع در حوزه‌های مختلف علمی بر روی وب جهان گستر قرار می‌گیرد. این مستندات می‌توانند شامل پایان‌نامه‌ها، مقاله‌ها، گزارش‌های علمی و مواردی از این قبیل باشد. نگارش متن این مستندات علمی جهت حفظ یکنواختی باید بر اساس اصول ثابت انجام گیرد، اما همواره به طور غیر عمدی دستخوش سلیقه‌های مختلفی در طول تاریخ می‌شود. اگرچه این تغییرات ناشی از پویا بودن زبان و خلاقیت ذهنی بشر است، اما این پویایی و خلاقیت پردازش ماشینی متن را با چالش‌های متعددی روبرو می‌کند و دقت پردازش داده‌ها را به میزان چشمگیری پایین می‌آورد. علاوه بر تنوع نگارشی، غلط‌های سهوی املائی نیز وجود دارد که فحوای گفتمانی متن را منحرف کرده و درک آن را با مشکل مواجه می‌کند. بنابراین کلیه نویسه‌های متن باید به حالت استاندارد تبدیل شوند و عاری از هرگونه خطاهای املائی گردند. پژوهشگران طرح حاضر سامانه‌ای برای استانداردسازی و خطایابی متون علمی فارسی طراحی کرده‌اند که این سامانه متون نوشتاری علمی و تخصصی فارسی را به لحاظ صحت نگارشی و املائی بررسی می‌کند و متن را به شکل استاندارد در می‌آورد. این سامانه با معیار اف ۹۰,۰۸ و معیار بلوی ۹۱,۶۶ به استانداردسازی می‌پردازد و با معیار اف ۹۲ درصد و با رتبه‌بندی ۸۲ درصد به اصلاح خطای املائی می‌پردازد.



زبان مهم‌ترین و کامل‌ترین وسیله نه تنها برای ایجاد ارتباط همزمان میان افراد، بلکه به منظور ارتباط بین نسل‌ها و فرهنگ‌های مختلف در گذشته، حال و آینده نیز به کار می‌رود. بنابراین نقش عمده زبان، ایجاد ارتباط و مراوده اجتماعی، چه در سطوح تاریخی و چه در سطوح جغرافیایی است، و انسان برای انتقال تجربیات و اندیشه‌های خود به دیگران به ابزاری که همان زبان است، نیازمند است. چنانچه رشد ملت‌های گوناگون را با گستردگی گنجینه لغات زبان آن‌ها مقایسه کنیم، درمی‌یابیم که بین نهادهای اجتماعی زبان و فرهنگ رابطه‌ای مستقیم وجود دارد. توسعه و گسترش افکار و اطلاعات مستلزم رشد متناسب گنجینه لغات است، و همچنین تاثیر زبان و تفکر را بر یکدیگر می‌رساند.



میزگرد و سخنرانی علمی

ایرانداک برای گسترش گفتمان اصطلاح‌شناسی و زبان‌شناسی رایانشی، در میزگردها و سخنرانی‌های علمی به زمینه‌های گوناگون این حوزه می‌پردازد و کارشناسان و استادان را برای گفت‌وگوهای انتقادی در این زمینه گرد می‌آورد و دستاوردهای آنها را در اختیار همگان می‌گذارد.

سخنرانی علمی: فعالیت‌های بین‌المللی اصطلاح‌شناسی و کاربرد آن در بازنمایی دانش

در این نشست که با حضور رئیس اینفوترم (International Information Centre for Terminology) که نهاد بین‌المللی فعالی در حوزه اصطلاح‌شناسی است برگزار شد، در ابتدا فعالیت‌های اصطلاح‌شناسی و اصطلاح‌نامه‌های علوم پایه تدوین یافته در گروه اصطلاح‌شناسی پژوهشگاه علوم و فناوری اطلاعات و آنتولوژی حاصل از آنها توسط اساتید ایرانداک معرفی شد. در ادامه، رئیس اینفوترم به تشریح نقش اصطلاح‌شناسی در سازماندهی و مدیریت دانش نتایج فعالیت‌های تحقیقاتی، نقش کلیدواژه‌ها و اصطلاحات در بازنمایی دانش و انتقال آن به متخصصین و پژوهشگران، کاربرد زبان و برنامه‌ریزی اصطلاح‌شناسی در سطح ملی و بین‌المللی، استانداردها، اصول و روش‌های اصطلاح‌شناسی، معرفی سازمان اینفوترم و نقش آن در هماهنگی فعالیت‌های اصطلاح‌شناسی در سطح بین‌الملل پرداختند.



سخنرانی علمی: مکتب اصطلاح‌شناسی

در این سخنرانی، ابتدا درباره سه معنای اصطلاح‌شناسی (ترمینولوژی) صحبت شد. این سه معنا شامل: ۱- اصطلاح‌شناسی به معنی مجموعه اصطلاحات علمی، ۲- اصطلاح‌شناسی به معنی فعالیت واژه‌گزینی و معادل‌یابی و ۳- اصطلاح‌شناسی به معنی رشته‌ای علمی. سپس، برخی از مکاتب

اصطلاح‌شناسی مانند مکتب وین، مکتب مسکو، مکتب پراگ، مکتب کبک، مکتب شناختی و مکتب اجتماعی معرفی شدند و در ادامه به تاثیرپذیری مکاتب اصطلاح‌شناسی از مکاتب زبان‌شناسی پرداخته شد. در نهایت به تفصیل درباره مکتب وین و تاثیرپذیری آن از برخی از مکاتب زبان‌شناسی توضیح داده شد و در این راستا به تقابل مکتب زبان‌شناسی با مکتب وین پرداخته شد.



سخنرانی علمی: تاکتیک‌های جست‌وجو و هستی‌شناسی‌ها

کاربران به صورت مداوم پرسش جست‌وجوی قبلی خود را به امید بازیابی بهتر اصلاح می‌کنند. این اصلاحات فرمول‌بندی مجدد یا اصلاح پرس‌وجو نامیده می‌شود. بیتس (۱۹۹۰) اشاره می‌کند، گاهی کاربر ترجیح می‌دهد که فرایند جست‌وجو را تا حدی، مطابق با استراتژی‌ها و تاکتیک‌های مشخصی که بتواند به طور موثر در رابط کاربر حمایت شود، کنترل کند. بیتس تاکتیک‌ها را به عنوان «یک یا تعداد انگشت‌شماری از حرکات که جست‌وجو را پیش می‌برند» تعریف می‌کند و به عنوان مثال، آنها را شامل تغییر یک اصطلاح به اصطلاحات خاص‌تر یا اضافه کردن اصطلاحات مشابه می‌داند. در این میان، هستی‌شناسی‌ها می‌توانند برای حمایت مستقیم تاکتیک‌ها استفاده شوند. (García and Sicilia 2003). در این سخنرانی، به تاکتیک‌هایی که از سوی هستی‌شناسی‌ها حمایت می‌شوند پرداخته شده و وضعیت کاربران گنج در استفاده از این تاکتیک‌ها به تصویر کشیده می‌شود.



سخنرانی علمی: رتبه‌بندی حوزه‌های پژوهشی زبان‌شناسی رایانشی با روش تحلیل سلسله‌مراتبی (AHP)

در این سخنرانی به مساله تعیین اولویت‌های پژوهشی گروه زبان‌شناسی رایانشی ایرانداک پرداخته شده و نتایج به دست آمده با استفاده از روش تحلیل سلسله‌مراتبی (AHP) برای این مساله ارائه شد. با توجه به مطالب ذکر شده، روش تحلیل سلسله‌مراتبی دارای دو سطح «گزینه‌ها و معیارها» است که در مساله موردنظر، گزینه‌های پژوهش حوزه‌های زبان‌شناسی رایانشی بوده که برای تعیین آنها از روش اصطلاح‌نامه‌ای استفاده شده، و معیارهای پژوهش بر اساس اسناد بالادستی پژوهشگاه و با نظر خبرگان بدست آمده است. علاوه‌براین، مشخص گردید که اولویت‌بندی معیارها و گزینه‌ها توسط خبرگان زبان‌شناسی صورت گرفته و در نهایت با تعیین وزن معیارها و گزینه‌ها، اولویت‌های پژوهشی بر اساس نظر خبرگان به دست آمده است.



سخنرانی علمی: رابطه خط و شخصیت

در این سخنرانی به رابطه بین دست‌خط افراد و شخصیت آنان پرداخته خواهد شد. از روی ویژگی‌های خط هر فرد و نحوه آرایش نوشتار و شیوه استفاده از قلم می‌توان به برخی ویژگی‌های

درونی افراد پی برد. شناسایی خلق و خوی نویسنده از روی دست خط به خصوصیات بومی خط نیز مربوط می شود.



سخنرانی علمی: از خواب نامه تا لغت نامه (ملاحظات در باب انتخاب واژگان بیگانه)

مشابهت غربی میان فرهنگ لغت و خواب نامه وجود دارد: هر لغت هم در میان لغت نامه وجود دارد و هم در خواب نامه. اما برای لغت های مدرن و نوظهور معادلی وجود ندارد. این نکته به خودی خود دیرهمضمی این لغات را در انبانه ناخود آگاه ذهن ایرانی نشان می دهد. هدف از این سخنرانی توجه دادن مخاطبان به رابطه خواب به عنوان تصویر (رویا= رویت) و کلمه و پرداختن به اهمیت انتخاب لغات جدید و تأثیر عمیقی است که می توانند بر ناخود آگاه بگذارند.



سخنرانی علمی: ایجاد هستی شناسی ها از اصطلاح نامه های موجود

در ابتدا ایجاد طرحواره مدیریت دانش در پژوهشگاه علوم و فناوری اطلاعات مطرح شد؛ طبق این طرحواره اسناد و مدارک به صورت دیداری و شنیداری، کتاب ها، پایان نامه ها و گزارش های دولتی

جمع‌آوری، سپس نمایه و وارد پایگاه‌های اطلاعاتی می‌شوند و نتایج تحلیل این اطلاعات علمی—فنی در اختیار سیاستگذاران کشور قرار می‌گیرند. مأموریت اصلی مرکز اطلاعات و مدارک علمی ایران در سالهای ۱۳۷۰ به عنوان یک مرکز اصلی و مادر در کشور، پشتیبانی مراکز تحقیقاتی و دانشگاهی و مراکز اجرایی و سیاست‌گذاری در زمینه اطلاعات علمی بود که این سازماندهی و مدیریت اطلاعات بدون کاربرد زبان‌های تخصصی و اصطلاح‌شناسی ممکن نبود. در این بخش اصطلاح‌شناسی و ارتباط آن با نشانه‌شناسی نمادهای کلامی و غیر کلامی، داده‌پردازی و انفورماتیک و مستندسازی، اطلاع‌رسانی و نظریه‌های علوم شناختی در این حوزه مطرح می‌شد. بر این اساس اصطلاح‌نامه به عنوان ابزار سازماندهی و بازیابی اطلاعات ایجاد گردید. در این سخنرانی فرایند شکل‌گیری اصطلاح‌نامه جامع، مراحل تشکیل هر اصطلاح‌نامه و فرایند تولید آن، مراحل ایجاد هستی‌شناسی‌ها با استفاده از اصطلاح‌نامه‌ها و مراحل ساخت ابزار مناسب برای انتقال اطلاعات یکسان شده از واژگان علوم پایه به نرم‌افزار آنتولوژی مورد بحث قرار گرفته است.

سخنرانی علمی: معرفی فارس‌نت: نخستین شبکه واژگانی زبان فارسی

با توجه به افزایش حجم اسناد و نیاز به پردازش ماشینی، عدم کارایی سیستم‌های پردازشی مبتنی بر نحو و کلیدواژه، نیاز روزافزون به سطوح بالاتر پردازش زبان (پردازش‌های معنایی)، نیاز به تهیه منابع زبانی و حرکت به سمت رویکردهای معنایی، ایجاد شبکه‌های واژگانی الزامی غیر قابل انکار است. در این سخنرانی، «واژه‌ستان‌شناسی» به عنوان یکی از گلوگاه‌ها در مسیر پردازش زبان فارسی بیان شده و برخی از کاربردهای آن مورد بررسی قرار گرفت. سپس، از «فارس‌نت» به عنوان پروژه‌ای عظیم برای ساخت واژه‌ستان‌شناسی در زبان فارسی نام برده شد که شبکه واژگانی (وردنت) فارسی، افزودن اطلاعات معنایی تکمیلی مانند قاب افعال، ساختارهای آرگومانی (نقش‌های موضوعی) و محدودیت‌های گزینشی و تبدیل واژگان معنایی به واژه‌ستان‌شناسی با افزودن روابط و ویژگی‌های متنوع‌تر، فازهای اول تا سوم آن را تشکیل می‌دهند. در ادامه، عناصر پایه‌ای وردنت فارسی معرفی شده و ساختار و ویژگی‌ها، گستره تحت پوشش، آمار کنونی دادگان، مفاهیم بنیادی، روابط میان مقوله‌ای و درون مقوله‌ای، روابط مشترک میان مقوله‌های نحوی، روابط خاص هر مقوله، روش‌های ایجاد شبکه‌های واژگانی، نحوه ساخت وردنت فارسی و ویژگی‌های آن، ابزارهای ساخت نیمه خودکار، نگاشت نیمه خودکار وردنت فارسی به انگلیسی، مشکلات نگاشت، و اشاره به برخی الگوها مانند الگوی هیرست و سایر الگوها مورد بحث واقع شد. در خاتمه این سخنرانی، به فعالیت‌های آتی در مسیر تکمیل فارس‌نت پرداخته شد.

سخنرانی علمی: ترجمه ماشینی مبتنی بر پیکره

یکی از کاربردهای مهم پیکره، استفاده از پیکره در ترجمه ماشینی است که نوعاً مبتنی بر استفاده از پیکره‌های موازی است. هدف از برگزاری این سخنرانی، معرفی نقش پیکره زبانی در توسعه و بهبود سامانه‌های ترجمه خودکار/ماشینی بود. در این راستا، ابتدا به تعریف مفهوم پیکره و ذکر برخی از انواع پیکره پرداخته شد. سپس، تاریخچه کوتاهی از صنعت ترجمه ارائه گردید. در ادامه به طور مفصل درباره معماری ترجمه و رویکردها و مدل‌های مطرح در آن مورد بررسی قرار گرفت و در نهایت، آخرین نوآوری‌ها و ویژگی‌های آن‌ها در ترجمه معرفی شد. نتایج بررسی انواع رویکردها با تکیه بر تاریخچه به کارگیری آن‌ها، نشان می‌دهد، مدل شبکه عصبی مبتنی بر پیکره، در مقایسه با سایر مدل‌ها (مدل مبتنی بر قاعده و مدل آماری مبتنی بر پیکره) تقریباً در همه جفت زبان‌هایی که بررسی شده است، عملکرد بهتری داشته است و خطاهای کمتری دیده شده است.



سخنرانی علمی: فرهنگ‌نگاری و اصطلاح‌نامه‌نگاری در زبان فارسی

در راستای مشخص نمودن وجه تمایز و تشابه «فرهنگ‌نگاری» با «اصطلاح‌نامه‌نگاری» در این رویداد، پس از بررسی پیکره‌های مهم‌ترین فرهنگ‌های فارسی به فارسی، از کهن‌ترین آن‌ها تا امروز، به معرفی طرح تدوین فرهنگ جامع زبان فارسی که زیر نظر دکتر علی‌اشرف صادقی در فرهنگستان زبان و ادب فارسی در دست اجراست، پرداخته شد و به بررسی مهم‌ترین رکن این فرهنگ یعنی پیکره رایانه‌ای آن، متشکل از شواهد به‌کارگیری واحدهای واژگانی در متون زبان فارسی پرداخته و ویژگی‌های برجسته آن شرح داده شد. همچنین در این جستار، به برخی از مهم‌ترین پیکره‌های زبانی رایانه‌ای که در سال‌های اخیر در ایران فراهم آمده نیز به اختصار اشاره گردید. سپس در سخنرانی دوم و سوم به تعریف اصطلاح‌نامه و روش تدوین و کاربردهای آن، پراخته شد. یکی از جدیدترین

کاربردهای اصطلاح‌نامه در امر آموزش است. در این راستا درباره کاربرد اصطلاح‌نامه در آموزش و تاثیر آن بر افزایش یادگیری فراگیران توضیح داده شد. در پایان، به تعریف مفهوم «فرااصطلاح‌نامه» و تفاوت آن با اصطلاح‌نامه پرداخته شد. سپس ویژگی‌ها و کاربرد فرااصطلاح‌نامه مطرح و در نهایت، نقش آن در ساخت آنتولوژی بررسی گردید.



سخنرانی علمی: نقش پیکره‌های زبانی در تحلیل گفتمان (تحلیل گفتمان پیکره‌یار)

هدف از برگزاری این رویداد معرفی مفهوم «پیکره»، «زبان‌شناسی پیکره‌ای»، کاربردهای پیکره با تمرکز روی یکی از کاربردهای مهم پیکره (تحلیل گفتمان) بود. بر این پایه، در این سخنرانی پس از تعریف «تحلیل گفتمان پیکره‌یار» ابتدا به بررسی وجه تمایز این حوزه را با پژوهش‌های پیکره‌بنیاد و پژوهش‌های برگرفته از پیکره، پرداخته شد. سپس، روش کار در تحلیل گفتمان پیکره‌یار توضیح داده شد. در ادامه یکی از نرم‌افزارهای پرکاربرد در تحلیل واژگانی متن معرفی شد. انواع کارایی این نرم افزار، «ورداسمیت»، شامل: قابلیت پردازش متون فارسی و انگلیسی، کاربرد در زمینه سبک‌شناسی، امکان جستجوی واژه‌ها و عبارت‌ها در متن، امکان انتخاب پیکره مورد بررسی برای جستجوی واژه‌ها و عبارت‌ها، مشخص کردن باهم‌آیندهای یک واژه و مشخص کردن الگوی توزیع یک واژه در متن است. با انتخاب گزینه‌های گوناگون این نرم‌افزار نتایج مختلفی از به کار رفتن واژه‌ها در بافت‌های گوناگون به دست می‌آید که می‌تواند تحلیل‌گر را در انواع تحلیل‌های گفتمانی از جمله مشخص نمودن سبک یک متن، نگرش و سوگیری خاص یک متن، تحلیل گفتمان متون سیاسی، تاریخی، مذهبی و.... یاری رساند.



سخنرانی علمی: زبان‌شناسی و کاربرد آن در حوزه بالینی (گفتاردرمانی)

هدف از برگزاری این رویداد، معرفی زبان‌شناسی بالینی به عنوان حوزه‌ای میان‌رشته‌ای و پرداختن به نقش دانش زبان‌شناسی در تشخیص و درمان بسیاری از اختلال‌های رشدی و اکتسابی زبان بود. محورهای این نشست شامل: پرداختن به مفهوم ارتباط و مولفه‌های آن، تعریف اختلال‌های رشدی زبان و عوامل پاتولوژیکی و غیرپاتولوژیکی مطرح، طبقه‌بندی انواع اختلال‌های رشدی زبان، نقش زبان-شناسی در پژوهش‌های گفتاردرمانی و اتخاذ رویکردهای درمانی مناسب، بررسی آسیب‌های زبان-شناختی در اختلال‌های زبانی و ارتباطی بزرگسالان، بررسی پارامترهای روان‌شناسی زبان، انواع درمان-ها در حوزه اختلال‌های اکتسابی زبانی و نقش زبان‌شناسی در انتخاب مولفه‌های گوناگون زبانی برای تحریک‌های زبانی مناسب است. جمع‌بندی مطالب ارائه شده در این نشست را می‌توان اینگونه خلاصه کرد: زبان‌شناسی در ارائه چارچوبی محکم و قابل استناد برای روند رشد طبیعی زبان در کودکان فارسی زبان، بررسی فرآیندهای گوناگون دخیل در اختلال‌های رشدی زبان، ساخت آزمون‌های گوناگون، بررسی ویژگی‌های زبان مربوط به هر کدام از اختلال‌ها و تهیه داده اولیه مربوط به روند رشد طبیعی زبان در کودک، می‌تواند به گفتاردرمانی کمک کند. در حوزه اختلال‌های زبانی مربوط به بزرگسالی نیز، انتخاب پارامترهای روان‌شناسی زبان و مشخصه‌های زبانی خاص به منظور تحریک‌های موثر زبانی، با همکاری زبان‌شناس صورت می‌پذیرد. همچنین در ساخت ابزارها و آزمون‌های خاص تشخیصی برای آلزایمر، دمانس و سایر اختلال‌های زبانی مربوط به بزرگسالان می‌توان به صورت تیمی (متشکل از گفتاردرمان، زبان‌شناس، روان‌شناس، عصب‌شناس، متخصص هوش مصنوعی و....) عمل کرد.



سخنرانی علمی: علوم انسانی در رایاسپهر

دوران کنونی را عصر انفجار اطلاعات، دوران موج سوم تحولات جهانی، انقلاب انفورماتیک و روزگار فناوری اطلاعات نامیده‌اند. ما در فضائی رایانشی و الکترونیکی (یعنی رایاسپهر) شناوریم. برای اینکه بتوانیم در چنین فضائی جایگاه شایسته خود را به دست آوریم و همدوش دیگران پیش رویم، چه باید بکنیم؟ وضعیت و جایگاه علوم انسانی بویژه زبان‌شناسی و زبان و فرهنگ ما چگونه است؟ نقش برنامه‌ریزان زبان، زبان‌شناسان و دست‌اندرکاران این حوزه چیست؟ برای پاسخ به این سوال‌ها، در این رویداد، ضمن ارائه توضیحات مفصل درباره مفهوم «رایاسپهر»، به بررسی جایگاه علوم و فنون گوناگون، به طور خاص علوم انسانی، در رایاسپهر پرداخته شد تا به این وسیله اهمیت امکانات موجود در رایاسپهر در پیشرفت‌های علوم انسانی مشخص گردد. سپس، به بررسی جایگاه زبان‌شناسی و زبان فارسی در رایاسپهر پرداخته شد و در نهایت گام‌های اساسی برای تثبیت بیشتر جایگاه خط و زبان فارسی در رایاسپهر معرفی گردید که این گام‌ها شامل: زمینه‌سازی، بسترسازی، ابزارسازی و محتواسازی است. اکنون سه گام اول در کشور به خوبی برداشته شده است و تنها چیزی که لازم داریم تولید محتوای مناسب صوری، فکری، علمی و... است.



سخنرانی علمی: پیکره و برجسب گذاری آن با کمک رایانه: ویژگی ها، کاستی ها و چالش ها

یکی از چالش های مطرح در زبان شناسی پیکره ای و رایانشی، نشانه گذاری داده ها (اضافه کردن اطلاعات زبانی توضیحی-تفسیری به پیکره) است. در این رویداد سعی شده است انواع نشانه گذاری (دستی و ماشینی) و ویژگی های مثبت و کم و کاستی های هر کدام از رویکردها معرفی شود. در این راستا برخی از مدل های ماشینی نشانه گذاری داده برای زبان فارسی، نوعاً درخت نحوی/تری بنک توضیح داده شده است. بررسی نقاط ضعف و قوت نحوه نشانه گذاری داده (دستی و ماشینی) نشان می دهد، یادگیری ماشینی فعال، روش مناسبی برای نشانه گذاری داده در حداقل حجم و کاهش زمان و هزینه است. همچنین در این رویداد به بررسی چالش های مربوط به داده های زبان فارسی، برجسب زنی و مدل یادگیرنده نیز پرداخته شده است. نتایج بررسی ها نشان می دهد علیرغم چالش ها، برجسب گذاری ماشینی بسیار بهتر از برجسب گذاری دستی است و این چالش ها در مقایسه با صرف زمان و هزینه برجسب گذاری دستی، ناچیز است.



سخنرانی علمی: اختصارسازی در زبان فارسی

در زبان فارسی کلمات و اصطلاحات تخصصی از رشته های مختلف وارد زبان شده است. این مسئله به خصوص در رشته هایی مانند پزشکی و رایانه که در زندگی مردم عادی بیشتر وارد شده اند، پر رنگ تر است. بنابراین مردم عادی مجبورند که واژه ها و اصطلاحات تخصصی همچنین سرنام های این حوزه های تخصصی را بشناسند و با آنها آشنا باشند تا مشکل تفهیم و تفاهم پیش نیاید. در این سخنرانی، پس از بیان مقدمات، دلایل پیدایش اختصارات بیان شد. سپس، صورت های مخمر، صورت های

ادغام شده و سرواژه‌ها به عنوان صورت‌های مختلف اختصارسازی بررسی شدند. در انتها نکات مورد نیاز به در نظر گرفته شدن در ایجاد صورت‌های سرواژه‌سازی ذکر شدند.

سخنرانی علمی: تحلیل عرضه و تقاضای پایان‌نامه‌ها و رساله‌ها (مطالعه موردی: حوزه انرژی)

این سخنرانی درصدد است راه کار مناسبی را مبتنی بر تحلیل محتوای پژوهش‌های موجود و تقاضاهای جامعه پژوهش در حوزه انرژی و از طریق کاربست تکنیک‌های متن کاوی و پردازش زبان طبیعی و تحلیل رفتار جویشگری کاربران ارائه نماید که از طریق آن کاستی‌های پژوهشی به صورت مستقل از نظر خبرگان و در فرایندی رایانشی شناسایی شود. محتوای این سخنرانی می‌تواند راهنمایی باشد تا سیاست‌گذاران و مدیران پژوهشی را در انتخاب روش تصمیم‌گیری کلان پژوهشی و اولویت‌بندی پژوهشی یاری رساند. همچنین به دلیل رویکرد میان‌رشته‌ای و روش تلفیقی که در این سخنرانی بیان خواهد شد، پژوهشگران و علاقه‌مندان حوزه‌های مختلف فناوری اطلاعات، هوش مصنوعی، کتابداری و اطلاع‌رسانی، سیاست‌گذاری علم و فناوری و ... می‌توانند از آن بهره‌مند شوند.

سخنرانی علمی: غنی‌سازی روابط معنایی در هستی‌شناسی شیمی

هستی‌شناسی، واژگان و مفاهیم مشترک مورد استفاده برای توصیف و ارائه یک حوزه از دانش را معین می‌کند و می‌تواند مبنایی برای سازماندهی و مدیریت دانش در یک دامنه موضوعی خاص باشد. به منظور تسهیل در این امر می‌توان از منابع دیگر دانش مانند واژه‌نامه‌ها و اصطلاح‌نامه‌ها استفاده کرد. اصطلاح‌نامه به لحاظ دارا بودن اطلاعات معنایی و ساختار سلسله‌مراتبی مفاهیم و نیز روابط معنایی میان آن‌ها منبع مناسبی برای ساخت هستی‌شناسی می‌باشد. به منظور ایجاد هستی‌شناسی از اصطلاح‌نامه شیمی که پیش از این در پژوهشگاه علوم و فناوری اطلاعات تدوین شده بود، استفاده شد و مبنای ساخت هستی‌شناسی شیمی قرار گرفت. طراحی مفهومی هستی‌شناسی با استفاده از روش «مت آنتولوژی» و بر اساس مفاهیم و روابط موجود در اصطلاح‌نامه صورت گرفت و اطلاعات آن از محیط اصطلاح‌نامه‌ای به نرم‌افزار پروتژه به طور خودکار انتقال داده شد و سپس روابط معنایی آن گسترش یافته و غنی‌سازی صورت گرفت.

سخنرانی علمی: ویژگی‌های نحوی زبان علم فارسی

زبان فارسی نیز همانند دیگر زبان‌ها به گونه‌های اجتماعی متفاوتی تقسیم می‌شود. این گونه‌ها تابع متغیرهای سن و جنسیت و حرفه و جایگاه اجتماعی گویشوران در جامعه زبانی است. زبان علم یکی از

گونه‌های اجتماعی زبان فارسی است که طبقات تحصیل کرده در مجامع و بافت‌های علمی آن را به کار می‌بندند. در این سخنرانی درباره ویژگی‌های صوری و زبانی مدارک علمی به ویژه استفاده از الفاظ فخیم فارسی و عربی در مقوله‌های اسم و صفت و قید و حروف اضافه و کاربرد فراوان ساخت مجهول و استفاده بیشتر از زمان افعال صحبت خواهد شد.

میزگرد تخصصی: زبان‌شناسی رایانشی

یکی از دغدغه‌هایی همیشگی ایرانداک به عنوان آرشیو ملی پایان‌نامه‌ها و رساله‌های دانش‌آموختگان تحصیلات تکمیلی کشور، ارتقا روش‌های ذخیره‌سازی، جستجو و بازیابی اطلاعات در زبان فارسی برای افزایش بهره‌وری و رضایت پژوهشگران کشور می‌باشد. با توجه به نقش زبان-شناسی رایانشی در ذخیره‌سازی، جستجو و بازیابی اطلاعات، ایرانداک زبان‌شناسی رایانشی را یکی از حوزه‌های پژوهشی اصلی خود قرار داده و علاوه بر ایجاد زیرساخت‌های فنی موردنیاز و تعریف و حل مسائل این حوزه توسط پژوهشگران پژوهشگاه، میزگردهایی تخصصی میان صاحب‌نظران این حوزه نیز برگزار کرده است.

ایرانداک سه میزگرد تخصصی با حدود ۱۰۰ نفر شرکت‌کننده در این زمینه برگزار کرده است. این میزگردها با عناوین زبان‌شناسی رایانشی، زبان‌شناسی رایانشی و علوم شناختی و پردازش زبان طبیعی و علوم شناختی برگزار شده است.



میزگرد تخصصی: اصطلاح‌شناسی در عصر حاضر

اصطلاح‌نامه گنجینه‌ای از واژه‌هاست که افزون بر نظم الفبایی فرهنگ‌ها، دارای نظامی شبکه‌ای و

مفهومی میان واژه‌های یک یا چند حوزه از دانش بشری است. این نظام مفهومی که ارتباطات گوناگون میان واژه‌ها را نشان می‌دهد، در سازمان‌دهی و بازیابی اطلاعات، تحلیل اطلاعات و علم‌سنجی، و آموزش کاربردهای بسیاری دارد. با توجه به اهمیت اصطلاح‌نامه‌ها، علاوه بر تدوین و به‌روزرسانی اصطلاح‌نامه‌های تخصصی به عنوان یک از اهداف پژوهشی اصلی در ایرانداک، این موضوع دست‌مایهٔ گفتمانی میان صاحب‌نظران این حوزه نیز قرار گرفته است.

ایرانداک میزگرد تخصصی در این زمینه برگزار کرده است که در آن سخنرانان درباره موضوعات «اصطلاح‌شناسی ابزاری برای سازمان‌دهی دانش»، «حرکت دیالکتیکی در پیدایش و خلق اصطلاح‌نامه‌ها و هستی‌شناسی‌ها»، «کاربردهای اصطلاح‌نامه در آموزش و تحلیل اطلاعات» و «کاربرد اصطلاح‌نامه در محیط دیجیتال» به ایراد سخنرانی پرداختند.



کارگاه و دوره آموزشی

ایرانداک آموزش‌های گوناگونی را برای گسترش چگونگی تدوین اصطلاح‌نامه‌ها و پردازش زبان طبیعی ارائه می‌دهد که شرکت در آنها برای همگان آزاد است. این آموزش‌ها تاکنون بارها برگزار شده‌اند که نام آنها در جدول زیر آمده است.

شمار	نام
۱۸	آشنایی با نمایه‌سازی و نمایه‌سازی ماشینی
۱۲	تدوین و کاربرد اصطلاح‌نامه
۸	آشنایی با شبکه‌های عصبی و استفاده از جعبه‌ابزار شبکه‌عصبی نرم‌افزار (MATLAB)
۴	پردازش زبان طبیعی
۴	آموزش تدوین اصطلاح‌نامه
۳	آشنایی با داده‌کاوی و کاربردهای آن
۲	وب معنایی: مفهوم‌سازی و کاربرد هستان‌شناسی (آنتولوژی)
۱	معرفی واژه‌شکن فارس
۱	فن ترجمه متون تخصصی

