

متن کاوی:

طبقه‌بندی، خوشه‌بندی و کاربردها

Text Mining: Classification, Clustering and Applications

تالیف:

سپنا دامی

(استادیار گروه مهندسی کامپیوتر، دانشگاه آزاد اسلامی واحد تهران غرب)

فرید سروری

(گروه مهندسی کامپیوتر، دانشگاه آزاد اسلامی واحد اردبیل)

عباس میرزایی

(استادیار گروه مهندسی کامپیوتر، دانشگاه آزاد اسلامی واحد اردبیل)



انتشارات جهاد دانشگاهی استان اردبیل

۱۴۰۰

سرشناسه	دلی، سینا، ۱۳۶۵ -
عنوان و نام پدیدآور	متن کاوی، طبقه‌بندی، خوشه‌بندی و کاربردها <i>Text mining, classification, clustering and -applications</i> تألیف سینا دلی، فرید سروری، عباس میرزایی
مشخصات نشر	اردبیل: جهاد دانشگاهی، واحد استان اردبیل، سازمان انتشارات، ۱۴۰۰
مشخصات ظاهری	۳۳۶ ص: ۴۵۰۰۰ ریال ۷-۸۱-۶۲۳۱-۶۲۲-۹۷۸
شابک	
وضعیت فهرست نویسی	فیا
موضوع	داده‌کاوری
موضوع	<i>Data mining</i>
موضوع	تجزیه و تحلیل کلاستر
موضوع	<i>Cluster analysis</i>
موضوع	تحلیل ممیز
موضوع	<i>Discriminant analysis</i>
موضوع	داده‌کاوری - روش‌های آماری
موضوع	<i>Data mining - Statistical methods</i>
شناسه افزوده	سروری، فرید، ۱۳۵۸ -
شناسه افزوده	میرزایی، عباس، ۱۳۶۱ -
شناسه افزوده	جهاد دانشگاهی، واحد اردبیل
شناسه افزوده	<i>Jahade Daneshgahi, Vahede Ardabil</i>
رده‌بندی کنگره	Q4 ۷۶/۹
رده‌بندی دیویی	۰۰۶/۳۱۲
شماره کتابشناسی ملی	۸۴۳۲۹۷۸

متن کاوی (طبقه‌بندی، خوشه‌بندی و کاربردها)



واحد استانی اردبیل

تألیف: سینا دلی، فرید سروری و عباس میرزایی

صفحه آرای: حسن قربانزاد

طرح جلد: محمود خروشی

ناشر: انتشارات جهاد دانشگاهی استان اردبیل

چاپ: اول - ۱۴۰۰

شمارگان: ۱۰۰۰ نسخه

قیمت: ۴۵۰۰۰ تومان

شابک: ۹۷۸-۶۲۲-۶۲۳۱-۸۱-۷

این اثر، مشمول قانون پشتیبانی مؤلفان و مصنفان و هنرمندان مصوب ۱۳۴۸ است. هر کسی تمام یا قسمتی از این اثر را بدون اجازه مؤلف (ناشر) نشر یا بخش یا عرضه نماید مورد پیگرد قانونی قرار خواهد گرفت.

نشانی انتشارات اردبیل - شهرک کارشناسان - میدان شفا، تلفن: ۰۳۲۷۴۹۵۰۶، راینامه: ard.chap@chmail.ir
نشانی سازمان انتشارات: تهران - خیابان انقلاب اسلامی - خیابان فخررازی - خیابان شهدای ژاندارمری - پلاک ۷۲، تلفن: ۶۶۹۵۲۹۴۸

پیشگفتار.....	۱۳
۱- مقدمه‌ای بر متن کاوی.....	۱۷
۱-۱ مقدمه.....	۱۸
۱-۲ الگوریتمها و کاربردهای متن کاوی.....	۲۰
۱-۳ جهت‌دهی‌های آینده.....	۲۴
۲- استخراج اطلاعات متن.....	۲۷
۲-۱ مقدمه.....	۲۸
۲-۲ شناسایی موجودیت نام‌دار.....	۳۲
۲-۲-۱ رویکرد مبتنی بر قاعده.....	۳۳
۲-۲-۲ رویکرد یادگیری آماری.....	۳۴
۲-۳ استخراج رابطه.....	۳۹
۲-۳-۱ طبقه‌بندی مبتنی بر ویژگی.....	۴۰
۲-۳-۲ روش‌های کرنل.....	۴۴
۲-۳-۳ روش‌های یادگیری نظارت ضعیف.....	۴۷
۲-۴ استخراج اطلاعات بدون نظارت.....	۴۸
۲-۴-۱ کشف رابطه و القاء الگو.....	۴۹
۲-۴-۲ استخراج اطلاعات آزاد.....	۵۰
۲-۵ ارزیابی.....	۵۱
۲-۶ خلاصه فصل.....	۵۲
۳- خلاصه‌سازی متن.....	۵۴
۳-۱ خلاصه‌سازی استخراجی.....	۵۵
۳-۲ رویکردهای بازنمایی موضوع.....	۵۷
۳-۲-۱ کلمات موضوع.....	۵۷

۱۶۳	۵-۶-۳ طبقه‌بندی شبکه عصبی
۱۶۶	۵-۶-۴ مشاهداتی درباره طبقه‌بندی خطی
۱۶۷	۵-۷ طبقه‌بندی مبتنی بر مجاورت
۱۷۰	۵-۸ طبقه‌بندی متن وب و داده‌های پیوندی
۱۷۶	۵-۹ فرا الگوریتم‌های طبقه‌بندی متن
۱۷۶	۵-۹-۱ طبقه‌بندی گروهی: یادگیری جمعی
۱۷۷	۵-۹-۲ طبقه‌بندی داده‌محور: بوستینگ و بگینگ
۱۷۹	۵-۹-۳ طبقه‌بندی دقت‌محور: بهینه‌سازی سنجش
۱۸۰	۵-۱۰ خلاصه فصل

۱۸۲	۶- متن کاوی در داده‌های جریانی
۱۸۳	۶-۱ مقدمه
۱۸۴	۶-۲ خوشه‌بندی جریان‌های متنی
۱۹۲	۶-۲-۱ شناسایی و ردیابی موضوع در جریان‌های متنی
۱۹۷	۶-۳ طبقه‌بندی جریان‌های متنی
۲۰۰	۶-۴ تحلیل تکامل در جریان‌های متنی
۲۰۲	۶-۵ خلاصه فصل

۲۰۳	۷- متن کاوی در داده‌های فرازبانی
۲۰۴	۷-۱ مقدمه
۲۰۵	۷-۲ ترجمه ماشینی
۲۰۵	۷-۲-۱ مدل‌های مولد ترجمه و SMT
۲۰۸	۷-۲-۲ مدل‌های مبتنی بر کلمه
۲۱۰	۷-۲-۳ مدل‌های مبتنی بر عبارت
۲۱۴	۷-۲-۴ مدل‌های مبتنی بر نحو
۲۱۷	۷-۳ کاوش متون موازی
۲۱۹	۷-۳-۱ کاوش ساختار وب
۲۲۱	۷-۳-۲ تطابق صفحات موازی
۲۲۴	۷-۴ مدل‌های ترجمه در CLIR
۲۲۶	۷-۵ جمع‌آوری و بهره‌برداری از متون تطبیقی

- ۲۳۰ ۷-۶ انتخاب کلمات ترجمه، عبارات و جملات موازی
- ۲۳۲ ۷-۷ کاوش روابط فرازبانی از متون تک‌زبانی
- ۲۳۵ ۷-۸ کاوش فرایوندها
- ۲۳۶ ۷-۹ خلاصه فصل

۸- متن کاوی در داده‌های چندرسانه‌ای ۲۳۹

- ۲۴۰ ۸-۱ مقدمه
- ۲۴۲ ۸-۲ کاوش متن پیرامون
- ۲۴۵ ۸-۳ برچسب کاوی
- ۲۴۵ ۸-۳-۱ رتبه‌بندی برچسب
- ۲۴۷ ۸-۳-۲ پالایش برچسب
- ۲۴۸ ۸-۳-۳ غنی‌سازی برچسب
- ۲۵۰ ۸-۴ کاوش محتوای مشترک متنی و بصری
- ۲۵۱ ۸-۴-۱ رتبه‌بندی بصری
- ۲۵۴ ۸-۵ کاوش محتوای متقابل متنی و بصری
- ۲۵۷ ۸-۶ خلاصه فصل

۹- متن کاوی در رسانه‌های اجتماعی ۲۶۰

- ۲۶۱ ۹-۱ مقدمه
- ۲۶۲ ۹-۲ جنبه‌های مجزای متن در رسانه‌های اجتماعی
- ۲۶۳ ۹-۲-۱ چارچوب عمومی برای تحلیل متن
- ۲۶۵ ۹-۲-۲ حساسیت به زمان
- ۲۶۶ ۹-۲-۳ طول کوتاه
- ۲۶۷ ۹-۲-۴ عبارات غیرساخت‌یافته
- ۲۶۸ ۹-۲-۵ اطلاعات اضافی
- ۲۶۸ ۹-۳ تحلیل متن در رسانه‌های اجتماعی
- ۲۶۹ ۹-۳-۱ شناسایی رویداد
- ۲۷۱ ۹-۳-۲ پاسخ‌گویی به سوال مشارکتی
- ۲۷۲ ۹-۳-۳ برچسب‌گذاری اجتماعی

۲۷۳	۹-۳-۴ پرکردن شکاف معنایی
۲۷۴	۹-۳-۵ بهره‌برداری از قدرت اطلاعات اضافی
۲۷۷	۹-۳-۶ تلاش‌های مرتبط
۲۷۷	۹-۴ مثال‌دنیای واقعی
۲۷۸	۹-۴-۱ استخراج عبارات اولیه
۲۸۰	۹-۴-۲ استخراج ویژگی‌های معنایی
۲۸۰	۹-۴-۳ ساخت فضای ویژگی
۲۸۲	۹-۵ خلاصه فصل

۱۰- عقیده کاوی و تحلیل احساسات

۲۸۶	۱۰-۱ مقدمه
۲۹۰	۱۰-۲ طبقه‌بندی احساسات
۲۹۲	۱۰-۲-۱ طبقه‌بندی احساسات مبتنی بر یادگیری با نظارت
۲۹۵	۱۰-۲-۲ طبقه‌بندی احساسات مبتنی بر یادگیری بدون نظارت
۲۹۸	۱۰-۳ طبقه‌بندی احساسات و اهمیت جمله
۲۹۹	۱۰-۴ گسترش واژه‌نامه عقیده
۳۰۰	۱۰-۴-۱ رویکرد مبتنی بر فرهنگ لغت
۳۰۰	۱۰-۴-۲ رویکرد مبتنی بر پیکره‌متنی
۳۰۲	۱۰-۵ تحلیل احساسات مبتنی بر جنبه
۳۰۳	۱۰-۵-۱ طبقه‌بندی جنبه احساسات
۳۰۵	۱۰-۵-۲ قواعد اساسی عقاید
۳۰۸	۱۰-۵-۳ استخراج جنبه
۳۱۰	۱۰-۵-۴ همزمانی گسترش فرهنگ لغت عقیده و استخراج جنبه
۳۱۱	۱۰-۶ عقیده کاوی تطبیقی
۳۱۵	۱۰-۷ برخی مسایل عقیده کاوی
۳۱۹	۱۰-۸ سودمندی عقاید
۳۲۰	۱۰-۹ خلاصه فصل

۱۱- متن کاوی معنایی مبتنی بر آنتولوژی

۳۲۳	۱۱-۱ مقدمه
-----	------------

- ۳۲۳ ۱۱-۲ وب معنایی و آنتولوژی
- ۳۲۶ ۱۱-۳ خوشه‌بندی معنایی متن
- ۳۲۷ ۱۱-۳-۱ تشابه معنایی
- ۳۲۹ ۱۱-۴ طبقه‌بندی معنایی متن
- ۳۳۱ ۱۱-۵ خلاصه فصل
- ۳۳۲ منابع

پیشگفتار

متن کاوی در سال‌های اخیر به دلیل حجم زیادی از داده‌های متنی که در شبکه‌های اجتماعی مختلف، وبسایت‌ها و سایر برنامه‌های کاربردی اطلاعات محور ایجاد شده، بیشتر مورد توجه قرار گرفته است. داده‌های غیرساخت‌یافته ساده‌ترین شکل داده‌هایی هستند که می‌توانند در هر سناریو ممکن ایجاد شوند. در نتیجه، یک نیازمیرم به طراحی روش‌ها و الگوریتم‌هایی احساس می‌شود که بتوانند طیف گسترده‌ای از برنامه‌های کاربردی مبتنی بر متن را پردازش کنند. این کتاب مروری بر روش‌ها و الگوریتم‌های مختلف و متداول در حوزه متن کاوی با تمرکز ویژه بر روش‌های طبقه‌بندی، خوشه‌بندی و کاربردهای آن خواهد داشت.

در ابتدا ضمن بیان مقدمه‌ای بر متن کاوی در فصل اول، به مفاهیم استخراج اطلاعات از متن و کارهایی که در جامعه پردازش زبان طبیعی انجام شده است، در فصل دوم می‌پردازیم. استخراج اطلاعات، کشف اطلاعات ساخت‌یافته از متن غیرساخت‌یافته نیمه‌ساخت‌یافته است. این یک کار مهم در استخراج متن است و به طور گسترده در مجامع مختلف پژوهش‌هایی از جمله پردازش زبان طبیعی، بازیابی اطلاعات و استخراج وب مورد مطالعه قرار گرفته‌است. دو وظیفه اساسی استخراج اطلاعات به عنوان شناسایی موجودیت‌های اسمی و استخراج روابط معنایی بین آن‌هاست است که در این فصل بدان پرداخته می‌شود.

وظیفه مهم دیگر، استخراج بخش‌های مهم از یک یا چند متن است که خلاصه‌سازی متن نام دارد. نوعی از خلاصه‌سازی متن را می‌توان به عنوان یکی از اهداف اصلی استخراج اطلاعات دانست. با توجه به این که اسناد زیادی در وب موجود است که بیشتر آن‌ها محتوای غیرضروری دارند، اهمیت خلاصه‌سازی متن به منظور کاهش زمان پردازش و مطالعه محسوس‌تر می‌شود. فصل سوم، به خلاصه‌سازی خودکار متن می‌پردازد که یکی از وظایف مهم و کاربردی در زمینه متن کاوی است. هدف از خلاصه‌سازی متن، یافتن خلاصه‌ای از محتوای متن به صورت فشرده و مطابق با اطلاعات مورد نیاز کاربر است. گونه‌ای که مفاهیم کلی و محتوای اطلاعاتی متن حفظ شود. تاکنون روش‌های بی‌شماری برای شناسایی محتوای مهم در خلاصه‌سازی متن توسعه یافته است. در این فصل، مروری کلی بر روی روش‌های موجود خلاصه‌سازی متن ارائه می‌دهیم. همچنین به برخی از مسائل و راه‌حل‌های پیشنهادی که رویکردهای یادگیری ماشین را برای این وظیفه به چالش کشیده‌است، اشاره می‌کنیم.

در فصل چهارم و فصل پنجم، به ترتیب به دو دسته‌بندی مهم متن کاوی، یعنی خوشه‌بندی متن و طبقه‌بندی متن می‌پردازیم. ما در فصل چهارم، یک بررسی جامع از مسأله‌ی خوشه‌بندی متن ارائه می‌کنیم. انواع روش‌ها، مزایای نسبی و چالش‌های کلیدی مسأله‌ی خوشه‌بندی بر روی داده‌های متن مورد مطالعه و بحث قرار می‌گیرند. در فصل پنجم، مروری جامع بر طیف گسترده‌ای از الگوریتم‌های طبقه‌بندی متن خواهیم داشت و به کاربردهای بی‌شماری نظیر طبقه‌بندی پرونده‌های

سلامت بیماران، نامه‌های الکترونیکی، مقالات علمی، متون خبری، صفحات وب و شبکه‌های اجتماعی اشاره می‌شود.

بخش اعظمی از داده‌های متنی که دائماً در طی زمان با کاربردهایی همچون شبکه‌های اجتماعی تولید می‌شوند، منجر به جریان‌های متنی‌کلان می‌گردد. عمدتاً این جریان‌های کلان متن از طریق تعاملات فردی در مقیاس عظیم یا از طریق تکوین ساختارمند انواع محتوای خاص توسط سازمان‌های اختصاصی ایجاد می‌شوند. بارزترین نمونه، جریان‌های متنی‌اخبار است که توسط سرویس‌های خبری تولید می‌شود. چنین جریان‌هایی، چالش‌های بی‌سابقه‌ای را برای الگوریتم‌های داده کاوی مطرح کرده است. ما در فصل ششم، الگوریتم‌های متن کاوی در داده‌های جریان‌یاب برای انواع مسائل موجود در داده کاوی از قبیل خوشه‌بندی، طبقه‌بندی و تحلیل تکامل بررسی می‌کنیم.

بازایی اطلاعات از میان چندین زبان مختلف مانند ترجمه متن کامل، عملی است که نیازمند نوعی انتقال دانش میان زبان‌ها می‌باشد. برای ساختن ابزارهای ترجمه باکیفیت وجود منابع ترجمه‌ی قوی امری حیاتی است. این در حالیست که منابع دست‌ساز دارای محدودیت‌هایی از قبیل شمول و سازگاری با طیف گسترده‌ای از کاربردها می‌باشند. متن کاوی در داده‌های فرازبانی ما را قادر می‌سازد تا این منابع دستی را تکمیل کنیم. فصل هفتم، به توصیف مباحثی می‌پردازد که در جست‌وجوی استخراج دانش فرازبانی از داده‌های متنی و به‌خصوص داده‌های موجود در وب هستند.

همان‌طور که می‌دانیم حجم بسیاری از داده‌های چندرسانه‌ای نظیر متن، تصویر و ویدیو در وب در دسترس است. معمولاً این رسانه‌ها با ترکیبی از انواع فراداده‌ها نظیر متن پیرامون، برچسب‌های کاربر، امتیازها، نظرات و غیره همراه می‌شود. کاوش این فراداده‌های متنی در تسهیل پردازش و مدیریت اطلاعات چندرسانه‌ای مؤثر است. فصل هشتم ارزیابی جامعی از متن کاوی در داده‌های چندرسانه‌ای فراهم می‌کند.

رشد سریع رسانه‌های اجتماعی آنلاین در قالب تولید محتوای مشارکتی، فرصت‌ها و چالش‌های جدیدی را برای تولیدکنندگان و مصرف‌کنندگان اطلاعات ایجاد می‌کند. با وجود تعداد زیادی از داده‌های تولید شده توسط سرویس‌های مختلف رسانه‌های اجتماعی، تحلیل متن، یک روش مؤثر برای پاسخگویی به نیازهای متنوع اطلاعاتی را فراهم می‌کند. در فصل نهم، ضمن معرفی پیشینه‌ی تحلیل متن و جنبه‌های متمایز داده‌های متنی در رسانه‌های اجتماعی، در مورد پیشرفت نتایج استفاده از متن کاوی در رسانه‌های اجتماعی از دیدگاه‌های مختلف بحث می‌کنیم.

یکی دیگر از جنبه‌های مهم رشد انفجاری رسانه‌های اجتماعی در وب، استفاده فزاینده‌ی افراد یا سازمان‌ها از نظرات و عقاید عمومی برگرفته از این رسانه‌ها برای تصمیم‌گیری است. کاوش متون مرتبط با عقاید و افکار مردم در وب به دلیل گسترش رسانه‌های متنوع، به عنوان یکی از کاربردهای

مهم متن کاوی مطرح است. هر وبسایت به طور معمول شامل حجم عظیمی از متن نظرات است که همیشه در پست‌ها و وبلاگ‌های انجمن‌های مرتبط رمزگشایی نمی‌شود. معمولاً یک خواننده‌ی انسانی در شناسایی وبسایت‌های مربوطه و خلاصه‌سازی دقیق اطلاعات و نظرات موجود در آن‌ها با مشکل روبه‌رو خواهد شد. علاوه بر این، همواره تحلیل انسانی از اطلاعات متنی در معرض تعصبات قابل توجهی قرار می‌گیرد؛ به عنوان مثال، افراد اغلب به عقایدیکه خودشان ترجیح می‌دهند تا با چیزی سازگارتر باشد، بیشتر توجه می‌کنند. آن‌ها همچنین به دلیل محدودیت‌های ذهنی و جسمی خود در پردازش حجم زیادی از اطلاعات با مشکل روبه‌رو می‌شوند که استفاده از سیستم‌های کاوش عقاید و خلاصه‌سازی نظرات را جهت تحلیل احساسات برای آن‌ها ملزم می‌کند. از این رو، فصل دهم را به مسائل مربوط به عقیده کاوی و تحلیل احساسات اختصاص دادیم.

در نهایت، پس از پرداختن به مسایل مربوط به متن کاوی در وب یک‌طرفه (وبسایت‌ها) و وب دوطرفه (رسانه‌های اجتماعی)، به استقبال مهم‌ترین مساله‌ی متن کاوی در نسل بعدی وب یعنی وب معنایی می‌رویم. هدف از وب معنایی این است که اطلاعات در وب نه تنها برای انسان‌ها قابل فهم باشد، بلکه ماشین‌ها نیز توانایی درک و پردازش آن‌ها را داشته باشند. یکی از مفاهیم اساسی وب معنایی، آنتولوژی است که در فصل یازدهم بعد از آشنایی با مفاهیم وب معنایی و آنتولوژی، مسائل مربوط به متن کاوی معنایی بر آنتولوژی مورد بررسی قرار می‌گیرد.

بدیهی است که هر اثری کامل و بدون نقص نیست از این رو، از خوانندگان علاقمند خواهشمندیم که نظرات و بازخورد سازنده‌ی خود را با ما از طریق پست الکترونیکی dami@wtiau.ac.ir به اشتراک بگذارند و در آینده نیز ما را همراهی کنند. در ادامه از تمامی دوستان و همکاران عزیزی که همراه و یاری‌گرمان بودند، صمیمانه سپاسگزاریم.

سینا دامی

بهار ۱۴۰۰