

روش‌های شناسایی داده‌های پرت

ترجمه: دکتر...

بهمن: فصل...



سازمان اسناد و کتابخانه ملی جمهوری اسلامی ایران

سرشناسه	: فرضی، بهناز، ۱۳۶۹
عنوان و نام پدیدآور	: روش‌های شناسایی داده‌ی پرت / نویسنده بهناز فرضی
مشخصات نشر	: تهران، آرون، ۱۳۹۹
مشخصات ظاهری	: ۹۲ ص.
شابک	: ۵ - ۹۳۰ - ۲۳۱ - ۹۶۴ - ۹۷۸
وضعیت فهرست‌نویسی	: فیپا
موضوع	: داده دورافتاده (آمار) - داده کاوی
موضوع	: Outliers (Statistics) - Data mining
رده‌بندی ده‌بندی	: QA ۲۷۶
رده‌بندی یوبی	: ۵۱۹/۵
شماره کتابخانه	: ۷۳۶۹۳۱۵



روش‌های شناسایی داده‌ی پرت

نویسنده: بهناز فرضی

ناشر: انتشارات آرون

چاپ اول: ۱۳۹۹

چاپ صدف: ۵۰۰

۲۷۰۰۰ تومان

نشانی: میدان انقلاب، خیابان ۱۲ فروردین، خیابان وحدت نظری

نرسیده به خیابان منیری جاوید، پلاک ۱۰۵، واحد ۳ تلفن: ۶۶۹۶۲۸۵۰ - ۶۶۹۶۲۸۵۱

ایمیل: Arvannashr@yahoo.com وبسایت: www.Arannashr.ir

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

فهرست مطالب

۷	چکیده
۹	فصل اول: کلیات کتاب
۹	مقدمه
۱۱	بیان مسئله
۱۲	پیشینه‌ی کتاب
۱۴	هدف کتاب
۱۴	اهمیت انجام کتاب
۱۶	معیار ارزیابی
۱۶	فرضیه‌های کتاب
۱۷	فرض‌ها و محدودیت‌های کتاب
۱۷	جنبه‌ی نوآوری کتاب
۱۸	ساختار کتاب
۱۹	فصل دوم: پیشینه‌ی کتاب
۱۹	مقدمه
۲۰	تکنیک امتیازدهی
۲۰	فاکتورهای شناسایی داده‌ی پرت
۲۲	طراحی سیستم تشخیص داده‌ی پرت
۲۴	کاربردهای شناسایی داده‌ی پرت
۲۵	تکنیک آماری

۲۵	تکنیک بر مبنای فاصله
۲۶	تکنیک بر مبنای چگالی
۲۷	تکنیک بر مبنای خوشه بندی
۲۸	تکنیک خوشه بندی بر مبنای تفکیک
۲۸	شناسایی داده‌های پرت توسط الگوریتم‌های خوشه بندی
۲۹	روش‌های خوشه بندی سلسله مراتبی
۳۳	روش دسته بندی بر اساس تعداد متغیر
۳۴	تکنیک بر مبنای همسایگی
۳۶	استفاده از روش‌های بهینه سازی
۳۸	الگوریتم‌های تکاملی
۴۰	به کارگیری الگوریتم ژنتیک
۴۲	به کارگیری الگوریتم بهینه سازی علف هرز مهاجم
۴۳	الگوریتم بهینه سازی رده ای مورچگان
۴۴	الگوریتم بهینه سازی گرابرات
۴۶	مدل بهینه سازی فارو
۴۷	الگوریتم تابکاری شبیه سازی داده
۴۸	الگوریتم جستجوی ممنوع
۴۹	الگوریتم اجتماع زنبور عسل
۵۱	تکنیک بر مبنای قانون
۵۱	تکنیک بر مبنای شبکه‌ی عصبی
۵۴	جمع‌بندی

فصل سوم: روش پیشنهادی برای شناسایی داده‌های پرت

۵۵	مقدمه
۵۵	رویکرد پیشنهادی این کتاب
۵۶	رویکرد پیشنهادی برای فاز اول کتاب: جستجوی تخمین نزدیک‌ترین همسایه
۵۸	رویکرد پیشنهادی برای فاز دوم کتاب: پیوندزنی روش‌ها
۵۹	الگوریتم تکاملی علف‌های هرز مهاجم
۶۲	روش خوشه بندی k-means
۶۴	روش حذف داده‌های پرت توسط خوشه بندی
۶۴	رویکرد پیشنهادی برای فاز سوم کتاب: حذف داده‌های پرت
۶۵	معیار ارزیابی الگوریتم
۶۶	عملکرد کلی الگوریتم
۶۸	جمع‌بندی

چکیده

در حال حاضر با افزایش روزافزون داده‌ها و حجم اطلاعات در مسائل دنیای پیرامون خود، روبرو هستیم که چالش روبروی ما، مدل سازی و تجزیه و تحلیل داده‌هاست و بهره گیری از روش‌هایی همچون داده کاوی برای استخراج دانش و اطلاعات نهفته در داده‌ها جزء مراحل ضروری شناسایی داده‌ها به شمار می‌آید. داده‌های پرت با عدم تطابق با سایر داده‌ها سبب بروز مشکل در امر تجزیه و تحلیل داده‌ها می‌گردند. بنابراین لزوم شناسایی داده‌های پرت امری اجتناب ناپذیر است. شناسایی داده‌های پرت و یا حصاها نقش مهمی در کاهش، محدود کردن حجم محاسبات دارند. داده‌های پرت در بسیاری از علوم کامپیوتری، پزشکی و تجارت کاربرد دارد. مسئله‌ی دیگری که امروزه در بحث داده کاوی وجود دارد، بحث کاهش خطا در شناسایی داده‌ها است. نقش شناسایی داده‌های پرت و کاهش خطاها، عامل اصلی مطالعه‌ی تکنیک‌های شناسایی داده‌های پرت و بهره‌برداری از تکنیک‌ها است.

در این کتاب سعی شده است بسیاری از تکنیک‌های شناسایی داده‌های پرت و معیارهای رده بندی این روش‌ها بیان گردد. از جمله این تکنیک‌ها می‌توان به تکنیک مبتنی بر خوشه بندی، تکنیک مبتنی بر همسایگی، تکنیک مبتنی بر چگالی و تکنیک امتیازدهی اشاره کرد. در این پایان نامه، ما قصد داریم تا ضمن مطالعه‌ی کارهای موجود در زمینه‌ی روش‌های شناسایی داده‌های پرت و کاربرد این روش‌ها در داده کاوی، روش جدیدی را با هدف کاهش خطا برای شناسایی این داده‌ها ارائه دهیم. ما مسئله خود را در سه فاز مورد مطالعه قرار خواهیم داد. در فاز اول، روش k نزدیک‌ترین همسایه مورد استفاده قرار می‌گیرد. در فاز دوم، به کمک الگوریتم علف‌های هرز، الگوریتم k -means اجرا می‌گردد. در فاز سوم، پس از تشکیل

خوشه‌ها با استفاده از روش k -means داده‌های پرت شناسایی شده، حذف می‌گردند. بطور کلی دستاوردهای اصلی این کتاب عبارتند از: (۱) ارائه‌ی روش اکتشافی نزدیک‌ترین هم‌سایه به منظور دست‌یابی به بهترین جواب در مسئله. (۲) ارائه‌ی یک روش پیوندی k -means و علف هرز به کمک روش نزدیک‌ترین هم‌سایه با هدف کاهش خطا در یافتن داده‌های پرت. نتایج حاصل از آزمایشات انجام شده در این پایان‌نامه، نشان‌دهنده برتری چشمگیری روش پیشنهادی با کاهش میانگین مربعات خطا در مجموعه‌ی داده‌ها می‌باشد.

www.ketab.ir