

۲۰۴۲۶۸۸

یادگیری ماشین تخصصی

نگرش امنیت اطلاعات

مؤلفین

آنتونی جوزف

بلین نلسون

بنیامین روبینستاین

ج.د. تایگار

مترجم

ایوب ترکیان

نیاز دانش

عنوان و نام پدیدآور	یادگیری ماشین تخصصی: نگرش امنیت اطلاعات/مؤلفین آنتونی جوزف... [و دیگران]؛ مترجم ایوب ترکیان.
مشخصات نشر	تهران: نیاز دانش، ۱۳۹۸.
مشخصات ظاهری	۳۹۲ ص: مصور، جدول، نمودار.
شابک	978-600-8906-50-6
وضعیت فهرست نویسی	فیفا
یادداشت	مؤلفین آنتونی جوزف- بلین، نلسون - بنیامین روبینستاین، ج.د. تایگار است.
یادداشت	عنوان اصلی: Adversarial machine learning, 2017
موضوع	Machine learning
موضوع	Computer security
موضوع	Joseph, Anthony D
شناسه افزوده	جوزف، آنتونی دی.
شناسه افزوده	ترکیان، ایوب، ۱۳۳۷، مترجم
رده بندی کنگ	Q۳۲۵ /۵
رده بندی دیویی	۰۰۶/۳۱
شماره کتابشناسی	۵۲۱۷۳۱۷



نام کتاب	یادگیری ماشین تخصصی نگرش امنیت اطلاعات
مؤلفین	آنتونی جوزف - بلین نلسون - بنیامین روبینستاین - ج.د. تایگار
مترجم	ایوب ترکیان
مدیر اجرایی - ناظر بر چاپ	حمیدرضا محمد شیرازی - محمد سمیع
ناشر	نیاز دانش
صفحه آرا	واحد تولید انتشارات نیاز دانش
نوبت چاپ	اول - ۱۳۹۸
شمارگان	۱۰۰ نسخه
قیمت	۶۸۰۰۰۰ ریال

ISBN:978-600-8906-50-6

شابک: ۹۷۸-۶۰۰-۸۹۰۶-۵۰-۶

هرگونه چاپ و تکثیر (اعم از زیراکس، پازنویسی، ضبط کامپیوتری و تهیهی CD) از محتویات این اثر بدون اجازه کتبی ناشر ممنوع است. متخلفان به موجب بند ۵ از ماده ۲ قانون حمایت از مؤلفان، مصنفان و هنرمندان تحت پیگرد قانونی قرار می گیرند.

کلیه حقوق این اثر برای ناشر محفوظ است.

آدرس انتشارات: تهران، میدان انقلاب، خیابان ۱۲ فروردین، تقاطع وحید نظری، پلاک ۲۵۵، طبقه ۱، واحد ۲

۰۲۱-۶۶۴۷۸۱۰۶-۶۶۴۷۸۱۰۸-۰۹۱۲۷۰۷۳۹۳۵

www.Niaze-Danesh.com

مشاوره جهت نشر: ۰۹۱۲-۲۱۰۶۷۰۹

فهرست مطالب

عنوان	شماره صفحه
فصل ۱ / مقدمه	۷
۱.۱ انگیزه	۹
۲.۱ رویکرد قاعده‌مند به یادگیری امن	۱۸
۳.۱ سیر زمانی یادگیری امن	۲۱
۴.۱ نمای کلان	۲۱
فصل ۲ / پس‌زمینه و نشانه‌گذاری	۷۳
۱.۱ نشانه‌گذاری پایه	۲۵
۲.۲ یادگیری ماشینی آماری	۲۶
۱.۲.۲ داده‌ها	۲۸
۲.۲.۲ فزاینده	۳۰
۳.۲.۲ مدل	۳۱
۴.۲.۲ یادگیری با نظارت	۳۱
۵.۲.۲ دیگر طرح‌واره‌های یادگیری	۳۵
فصل ۳ / چارچوب یادگیری امن	۱۱۳
۱.۳ تحلیل فازهای یادگیری	۳۸
۲.۳ تحلیل امنیت	۴۰
۱.۲.۳ اهداف امنیت	۴۰
۲.۲.۳ مدل تهدید	۴۱
۳.۲.۳ کاربردهای یادگیری ماشین در امنیت	۴۲
۳.۳ چارچوب	۴۳
۱.۳.۳ گروه‌بندی	۴۳
۲.۳.۳ بازی یادگیری تخصصی	۴۵
۳.۳.۳ خصوصیات توانمندی‌های تخصصی	۴۶
۴.۳.۳ حملات	۴۸
۵.۳.۳ دفاع‌ها	۴۹
۴.۳ حملات کاوشی	۴۹
۱.۴.۳ بازی کاوشی	۵۰
۲.۴.۳ حملات انسجام کاوشی	۵۱
۳.۴.۳ حملات قابلیت دسترسی کاوشی	۵۶
۴.۴.۳ دفاع مقابل حملات کاوشی	۵۷
۵.۳ حملات علی	۶۳
۱.۵.۳ بازی علی	۶۴
۲.۵.۳ حملات انسجام علی	۶۵
۳.۵.۳ حملات در دسترس بودن علی	۶۷
۴.۵.۳ دفاع مقابل حملات علی	۶۹
۶.۳ بازی‌های یادگیری تکرار شده	۷۳

۷۶	۱۶.۳ بازی‌های یادگیری تکرار شده در امنیت
۷۸	۷.۳ یادگیری سیانت حریم خصوصی
۷۸	۱.۷.۳ حریم خصوصی دیفرانسیل
۸۰	۲.۷.۳ حملات حریم خصوصی کاوشی و علی
۸۱	۳.۷.۳ منفعت علیرغم راندوم بودن

فصل ۴ / حمله به یادگیرنده فراکره‌ای

۸۵	۱.۴ حملات علی به آشکارسازهای فراکره‌ای
۸۶	۱.۱.۴ مفروضات یادگیری
۸۷	۲.۱.۴ مفروضات مهاجم
۸۸	۳.۱.۴ روش‌شناسی تحلیلی
۸۹	۲.۴ توصیف حمله فراکره‌ای
۹۲	۲.۴ جابه‌جا کردن مرکز ثقل
۹۶	۲.۱.۴ ضعیف ساختارمند حمله
۹۸	۲.۴ ویژگی توالی‌های حمله
۱۰۲	۳.۴ حملات نامقدّم بهینه
۱۰۳	۱.۳.۴ پشته ردن بلندها
۱۰۴	۴.۴ تحمیل قیود زمانی روی حمله
۱۰۶	۱.۴.۴ پشته کردن بهینه‌ها با جرم متغیر
۱۰۸	۲.۴.۴ فرمولاسیون اثرات
۱۰۹	۳.۴.۴ راه‌حل رهاشده بهینه
۱۱۳	۵.۴ حمله علیه بازآموزی با جایگزینی داده
۱۱۵	۱.۵.۴ سیاست جایگزینی میانگین آفر راندوم حذفی
۱۱۷	۲.۵.۴ سیاست جایگزینی نزدیک حذفی
۱۲۱	۶.۴ مهاجمین مقید
۱۲۳	۱.۶.۴ حملات بهینه مقتصدانه
۱۲۴	۲.۶.۴ حملات با داده‌های مخلوط
۱۲۸	۳.۶.۴ بسط‌ها
۱۲۹	۷.۴ خلاصه

فصل ۵ / مطالعه موردی حمله موجود بودن: SpamBayes

۱۳۳	۱.۵ فیلتر اسپم SpamBayes
۱۳۳	۱.۱.۵ الگوریتم آموزش
۱۳۵	۲.۱.۵ پیش‌بینی‌های SB
۱۳۶	۳.۱.۵ مدل SB
۱۴۱	۲.۵ مدل تهدید برای SpamBayes
۱۴۱	۱.۲.۵ اهداف مهاجم
۱۴۲	۲.۲.۵ دانش مهاجم
۱۴۳	۳.۲.۵ مدل آموزش دادن
۱۴۴	۴.۲.۵ فرض آلاش
۱۴۵	۳.۵ حملات علی علیه یادگیرنده SB
۱۴۵	۱.۳.۵ حملات موجود بودن علی
۱۵۰	۲.۳.۵ حملات انسجام علی - شبه‌اسپم
۱۵۰	۴.۵ دفاع ريجکت روی اثر منفی (RONI)

۱۵۲		SpamBayes با آزمایشات
۱۵۲		۱.۵.۵ روش تجربی
۱۵۴		۲.۵.۵ نتایج حمله فرهنگ‌نامه
۱۵۶		۳.۵.۵ نتایج حمله متمرکز
۱۶۰		۴.۵.۵ آزمایشات حمله شبه‌اسپم
۱۶۲		۵.۵.۵ نتایج RONI
۱۶۴		۶.۵ خلاصه

فصل ۶ / مطالعه موردی حمله انسجام: آشکارساز PCA ۲۹۳

۱۷۳		۱.۶ روش PCA آشکارسازی ناهنجاری‌های ترافیکی
۱۷۳		۱.۱.۶ ماتریس‌های ترافیک و ناهنجاری‌های حجم
۱۷۵		۲.۱.۶ روش زیرفضا برای آشکارسازی ناهنجاری
۱۷۷		۳.۱.۶ ایجاد اختلال در زیرفضای PCA
۱۷۷		۱.۱.۶ مدل تهدید
۱۷۹		۲.۲.۶ انتخاب آتشاک ناآگاهانه
۱۷۹		۳.۲.۶ انتخاب آگاهانه محلی
۱۷۹		۴.۲.۶ انتخاب آگاهانه فراگیر
۱۸۱		۵.۲.۶ حملات قه‌ق‌جوان
۱۸۳		۳.۶ آشکارسازهای ناب‌آورد مقابل اختلال
۱۸۳		۱.۳.۶ حس درونی
۱۸۵		۲.۳.۶ PCA-GRID
۱۸۷		۳.۳.۶ حد‌آستانه لاپلاس مستقیم
۱۹۰		۴.۶ ارزیابی تجربی
۱۹۰		۱.۴.۶ برپایش
۱۹۲		۲.۴.۶ شناسایی جریان‌های آسیب‌پذیر
۱۹۵		۳.۴.۶ ارزیابی حملات
۱۹۷		۴.۴.۶ ANTIDOTE ارزیابی
۱۹۹		۵.۴.۶ ارزیابی تجربی حمله مسموم‌سازی قورباغه جوشان
۲۰۵		۵.۶ خلاصه

فصل ۷ / مکانیسم‌های حفظ حریم خصوصی یادگیری SVM ۳۴۳

۲۰۷		۱.۷ مطالعات موردی رخنه به حریم
۲۰۸		۱.۱.۷ سوابق سلامت کارکنان دولتی
۲۰۸		۲.۱.۷ آگ‌های پرس‌مان موتور جستجوی AOL
۲۰۹		۳.۱.۷ جایزه Netflix
۲۰۹		۴.۱.۷ کشف نام شبه‌نام‌های تویتر
۲۱۰		۵.۱.۷ مطالعات انجمنی سطح ژنوم
۲۱۰		۶.۱.۷ هدف‌گذاری تبلیغات میکرو
۲۱۱		۷.۱.۷ آموخته‌ها
۲۱۲		۲.۷ موقعیت مسئله: یادگیری حفظ حریم
۲۱۲		۱.۲.۷ حریم دیفرانسیل
۲۱۵		۲.۲.۷ منفعت
۲۱۶		۳.۲.۷ سوگیری تاریخی پژوهش در حریم دیفرانسیل
۲۱۹		۳.۷ ماشین‌های بردار پشتیبان: معرفی اجمالی

۲۳۱	۱.۳.۷ کرنل های تغییرناپذیر به جابه جایی
۲۳۱	۲.۳.۷ پایداری الگوریتمی
۲۳۲	۴.۷ حریم دیفرانسیل برحسب اختلال خروجی
۲۳۷	۵.۷ حریم دیفرانسیل برحسب اختلال هدف
۲۳۰	۶.۷ فضاها و ویژگی ابعاد نامعین
۲۳۹	۷.۷ مرزهای حریم دیفرانسیل بهینه
۲۳۹	۱.۷.۷ مرزهای بالایی
۲۴۲	۲.۷.۷ مرزهای پایینی
۲۴۵	۸.۷ خلاصه

۳۴۳	فصل ۸ / گریز نزدیک بهینه طبقه گرها
۲۵۱	۱.۸ توصیه گریز نزدیک بهینه
۲۵۲	۱.۸ هزینه تخصصی
۲۵۴	۲.۱.۸ گریز نزدیک بهینه
۲۵۶	۳.۸ اصطلاحات جستجو
۲۵۹	۴.۱.۸ بهینه سازی ضرب جمع زدنی
۲۶۲	۵.۱.۸ خانواده طبقه گرها و القاگر تحذب
۲۶۴	۲.۸ گریز طبقات محذب و مرزهای ℓ_1
۲۶۵	۱.۲.۸ جستجوی $IMAC - \ell_1$ بر $\mathcal{X}_f^+$ محذب
۲۷۷	۲.۲.۸ یادگیری $IMAC - \ell_1$ برای $\mathcal{X}_f^+$ محذب
۲۸۵	۳.۸ گریز برای هزینه های ℓ_p کلی
۲۸۵	۱.۳.۸ مجموعه مثبت محذب
۲۹۲	۲.۳.۸ مجموعه منفی محذب
۲۹۳	۴.۸ خلاصه
۲۹۳	۱.۴.۸ مسایل باز در گریز نزدیک بهینه
۲۹۶	۲.۴.۸ معیارهای گریز اترناتیو
۲۹۹	۳.۴.۸ گریز دنیای واقعی

۳۴۳	فصل ۹ / چالش های یادگیری ماشین تخصصی
۳۰۹	۱.۹ بحث و مسایل باز
۳۰۹	۱.۱.۹ مؤلفه های بررسی نشده بازی تخصصی و مسایل باز
۳۱۱	۲.۱.۹ توسعه فناوری های دفاعی
۳۱۴	۲.۹ نکات پایانی
۳۱۷	پیوست الف/ پس زمینه یادگیری و فراهنده
۳۲۹	پیوست ب/ اثبات کامل حملات فرا کره
۳۳۹	پیوست ج/ اثبات SpamBayes فصل ۵
۳۴۹	پیوست د/ اثبات کامل گریز نزدیک بهینه
۳۵۹	پیوست ه/ متخاصم مبین