



آزمایشگاه شبکه‌های اطلاعاتی و کامپیوتری
دانشکده علوم و فنون نوین، دانشگاه تهران

تشخیص سرقت علمی با رویکرد مبتنی بر گراف

ارائه دهنده: مژگان ممتاز (m.momtaz92@ut.ac.ir)

ارائه مربوط به همایش بررسی راهکارهای پیشگیری از سرقت علمی

سایر نویسندگان: دکتر مصطفی صالحی، دکتر هادی ویسی، کیوان بیجاری

بهمن ۹۵

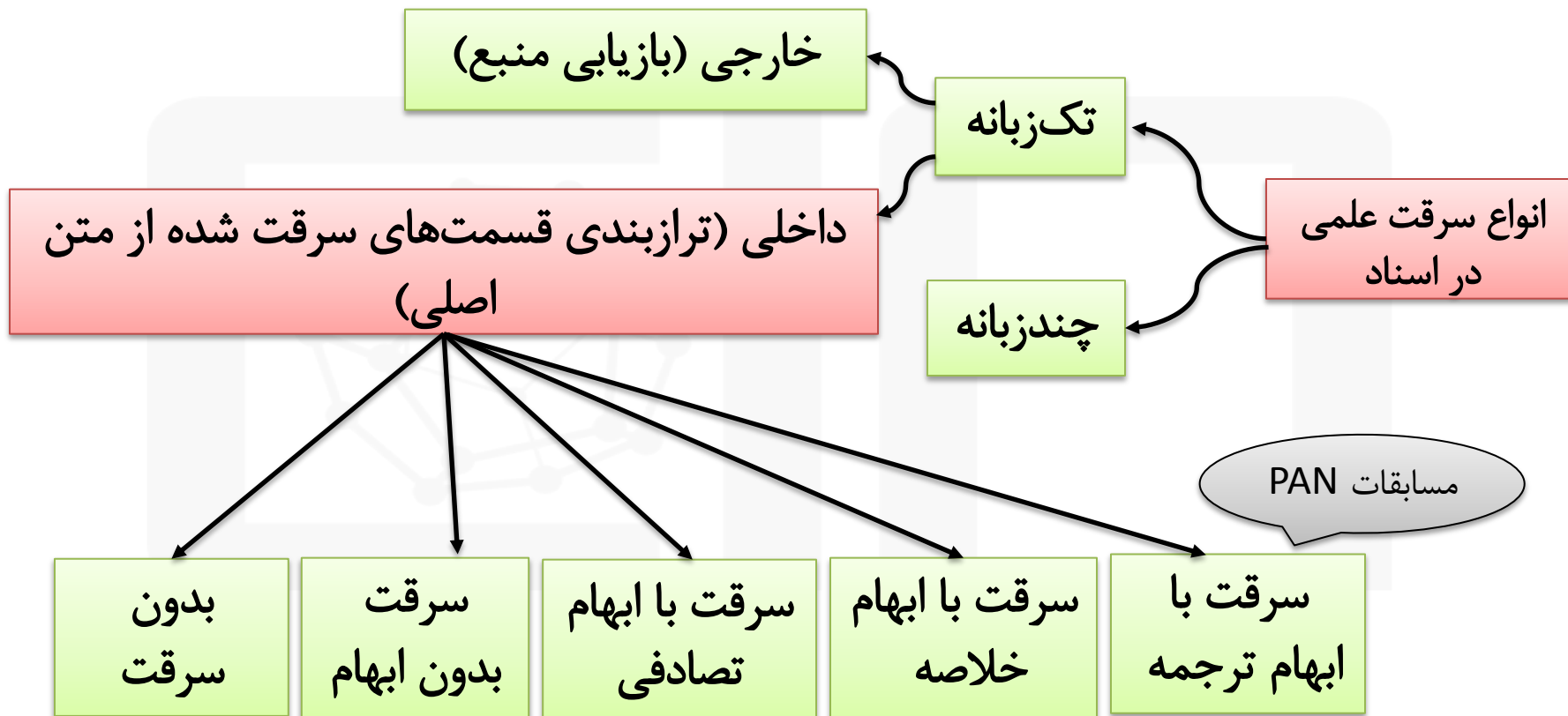
- ① مقدمه
- ① مروری بر ادبیات موضوع
- ① روش پیشنهادی
- ① ارزیابی
- ① نتیجه‌گیری و کارهای آتی
- ① مراجع

تعریف موضوع - سرقت علمی



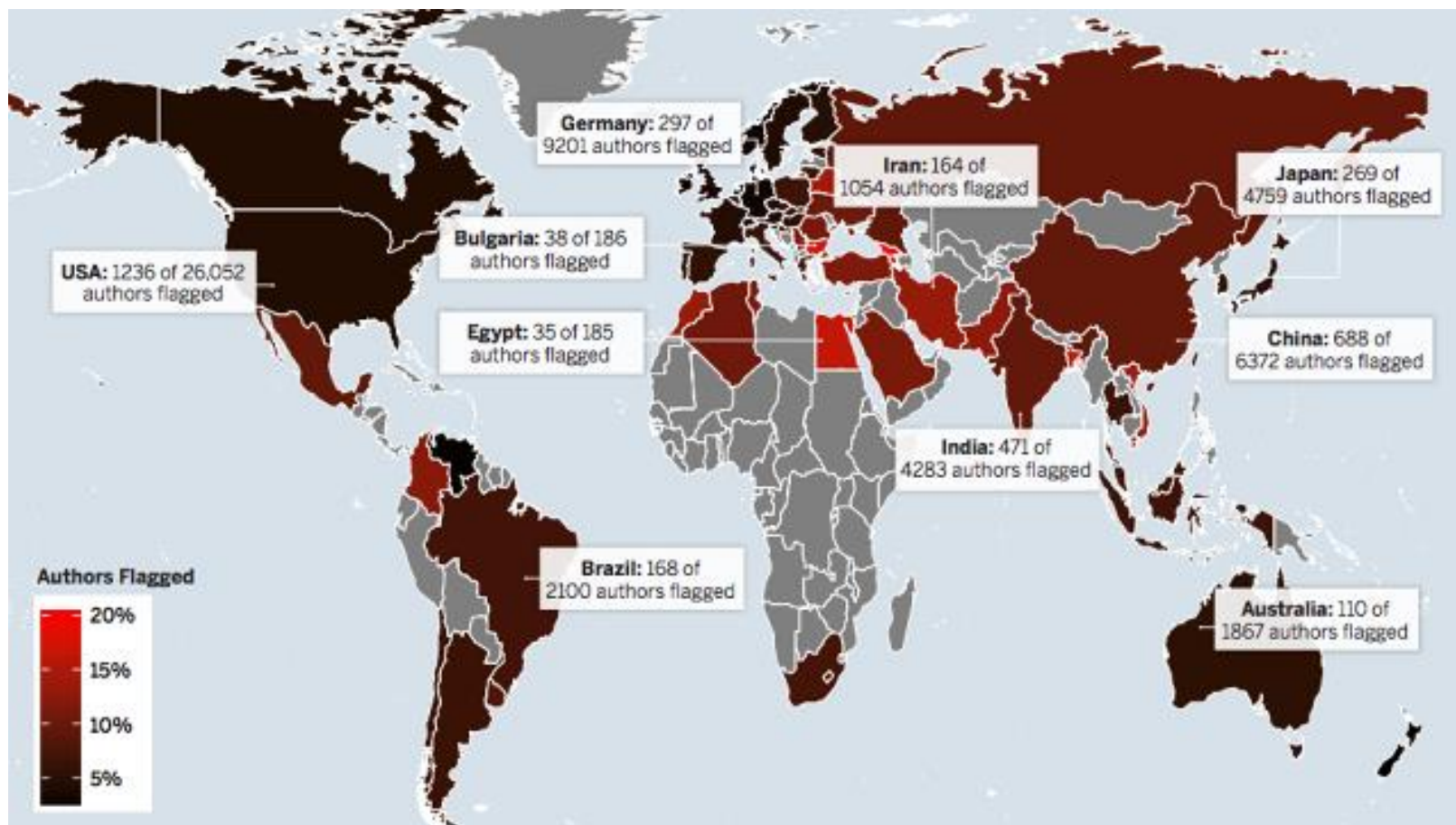
سرقت علمی - ادبی استفاده از ایده‌ها و عبارات دیگران، به عنوان ایده و عبارات خویش و در واقع نوعی دستبرد فکری یا ایده دزدی است. سرقت علمی می‌تواند شامل طیف وسیعی از دستبردهای آگاهانه تا کپی کردن اتفاقی مطالب دیگران باشد.

تعریف موضوع-انواع سرقت علمی در اسناد



⊛ بر اساس گزارش وبسایت نشریه Science، همه کشورهای مورد مطالعه کمابیش درگیر مسئله سرقت علمی در مجامع علمی هستند ولی این پدیده در برخی از کشورها بیشتر رخ داده است.

⊛ در این تحقیق گزارش شده است که بیش از پانزده درصد پژوهشگران ایرانی که در مقالات خود سرقت علمی داشته‌اند.



شکل ۱- نقشه سرقت علمی در جهان (sciencemag.org)



روش‌های پیشین تشخیص سرقت علمی

شماره	روش
۱	پیدا کردن مشابهت به کمک روش انگشتنگاری
۲	پیدا کردن مشابهت با روش n -گرام
۳	روش مبتنی بر خوشه‌بندی (بازیابی منبع)
۴	پیدا کردن مشابهت با روش ساختاری
۵	روش‌های مبتنی بر شباهت معنایی
۶	روش مبتنی بر گراف



مشکلات روش‌های پیشین نمایش متن

مشکلات روش‌های پیشین	رفع مشکلات در روش پیشنهادی
بی‌توجهی به ترتیب کلمات	استفاده از گراف جهت‌دار
بی‌توجهی به ساختار متن	توجه به کل جمله
بی‌توجهی به کلمات مترادف استفاده شده	امکان اضافه کردن کلمات مترادف به گراف متناظر متن
دقت پایین در تشخیص سرقت علمی از نوع خلاصه	تشخیص بهتر سرقت علمی نوع خلاصه با توجه به ویژگی انتخاب تعداد گره‌های همسایه‌های بیشتر

متدلوژی و روش – نمودار مراحل تشخیص سرقت علمی سطح بازیابی منبع

بسته‌های استفاده شده در پیاده سازی

روش پیشنهادی

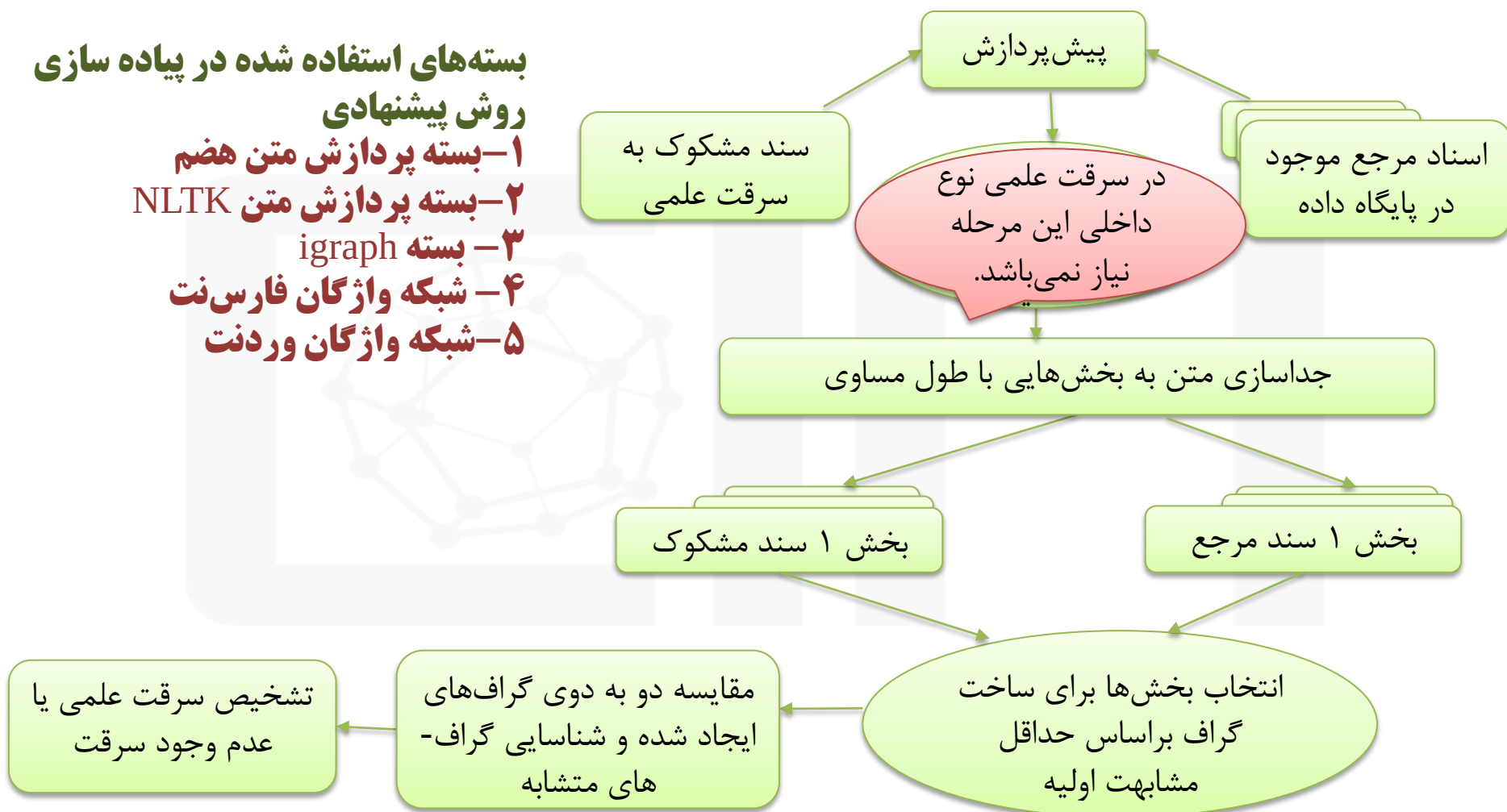
۱- بسته پردازش متن هضم

۲- بسته پردازش متن NLTK

۳- بسته igraph

۴- شبکه واژگان فارسی نت

۵- شبکه واژگان وردنت



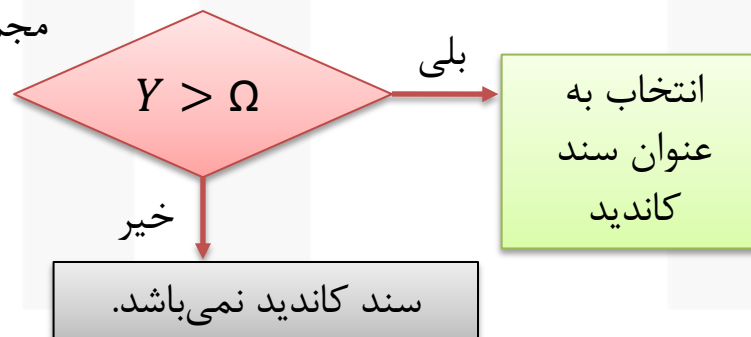
متدلوژی و روش-۱

مرحله پیش‌پردازش

- ⊛ استفاده از بسته هضم در نرمال‌سازی متن
- ⊛ برچسب‌زنی جز کلام (برای مرحله انتخاب سندهای کاندید)

مرحله انتخاب سندهای کاندید (سرقت علمی نوع خارجی)

- ⊛ مجموعه دوتایی‌ها سند مشکوک به سرقت علمی: $S1$
- ⊛ مجموعه دوتایی‌ها سند مرجع: $S2$
- ⊛ $Y = |S1 \cap S2|$ = تعداد دوتایی‌های مشترک
- ⊛ Ω = آستانه تعداد دوتایی‌های مشترک



- ⊛ نکته: دوتایی‌ها تنها شامل کلمات با تگ اسم می‌باشند و سایر تگ‌ها مانند صفت‌ها و قیدها حذف شده‌اند.

متدلوژی و روش-۲

⊛ مرحله جداسازی متن

- ⊛ جداسازی متن به مجموعه‌ای از پاراگراف‌ها ----- تشخیص سرقت علمی با ابهام خلاصه
- ⊛ جداسازی متن به مجموعه‌ای از جملات ----- تشخیص سرقت علمی سایر سطوح داخلی
- ⊛ جداسازی متن به بخش‌هایی با طول مشخصی از کاراکترها
- ⊛ جداسازی متن به بخش‌هایی با تعداد مشخصی از کلمات ----- تشخیص سرقت علمی سطح خارجی

⊛ معیار مشابهت اولیه جملات

- ⊛ شباهت کسینوسی

متدلوژی و روش - ۳

مرحله ساخت گراف برای بخش‌هایی با حداقل مشابهت اولیه

گروه: سند، مفهوم، پاراگراف، جمله، گروه اسمی، کلمه

برقراری یال: دستوری، معنایی، آماری، رخداد همزمان کلمات

✓ جمله

✓ پاراگراف

✓ متن

✓ پنجره با اندازه ثابت

در روش پیشنهادی:

گروه: کلمات (به جز ایست‌واژه‌ها)

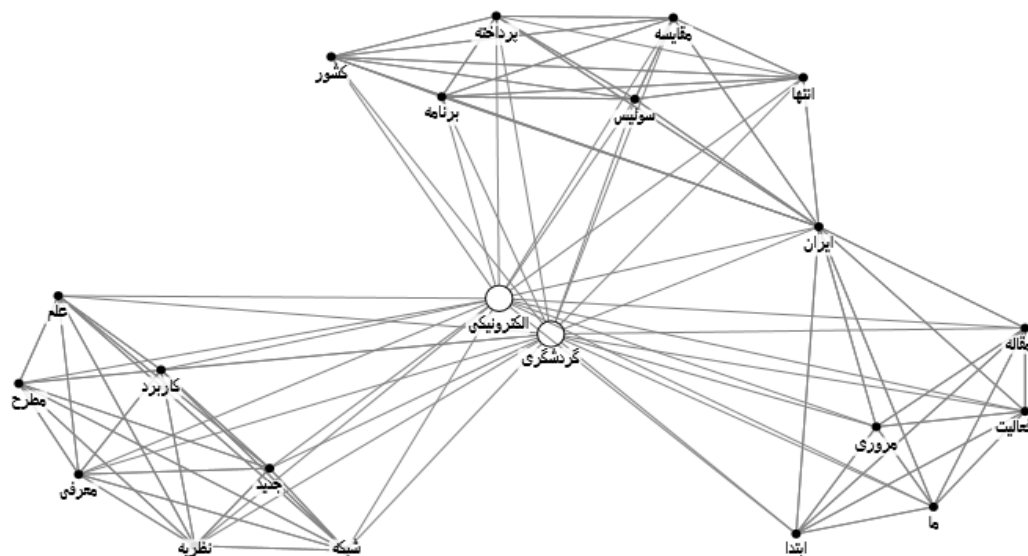
برقراری یال: رخداد همزمان

کلمات در پنجره با اندازه ثابت

مرحله ساخت گراف برای بخش‌هایی با حداقل مشابهت اولیه-ادامه

⊛ نمونه‌ای از تبدیل متن به گراف

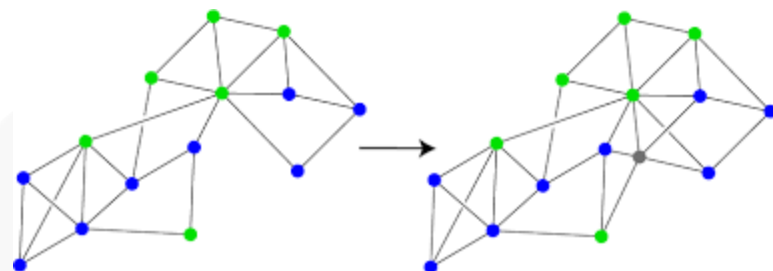
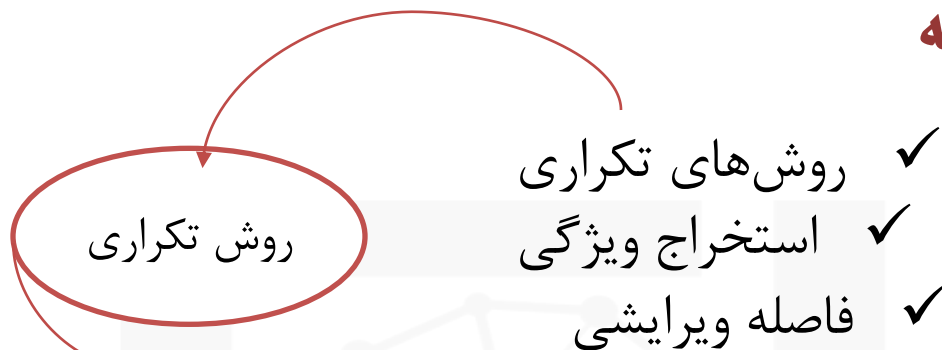
⊛ در این مقاله ما در ابتدا به مروری بر فعالیت‌های گردشگری الکترونیکی در ایران پرداخته‌ایم و سپس مقایسه‌ای بین برنامه‌های گردشگری الکترونیکی ایران با کشور سوئیس ارائه شده است، و در انتها پس از معرفی علم جدید نظریه شبکه‌ها، کاربرد آن را در گردشگری الکترونیکی مطرح می‌کنیم.



✓ رخداد همزمان کلمات در جملات

مدل‌وژی و روش-۵

مرحله تشخیص گراف‌های مشابه



⊛ همسایگان درجه اول گره مشترک گراف سند مشکوک A:

⊛ همسایگان درجه اول گره مشترک گراف سند مرجع B:

$$\star \text{similarity}(A, B) = \frac{A \cap B}{\max(\text{len}(A), \text{len}(B))}$$

آستانه مشابهت گره‌ها: α

if ($\text{similarity}(A, B) > \alpha$) then A similar to B

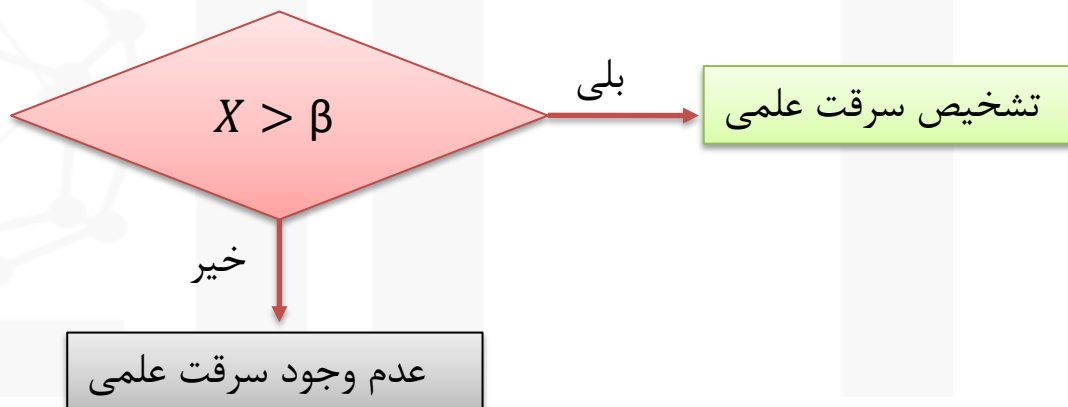
متدلوژی و روش - ۶

مرحله تشخیص گراف‌های مشابه - ادامه

آستانه تعداد گره‌ها مشابه مشخص می‌کند آیا دو گراف مشابه می‌باشند یا خیر؟

تعداد گره‌های مشابه: X

آستانه تعداد گره‌ها مشابه: β



ارزیابی - مجموعه داده فارسی (بازیابی منبع)

☆ مجموعه داده سرقت علمی معنایی و با درجه ابهام بالا

تعداد سند مرجع	تعداد سند مشکوک
۱۵۰۰	۱۱۲

☆ مجموعه داده همه سطوح سرقت علمی (معنایی، با ابهام، بدون ابهام، بدون سرقت علمی)

تعداد سند مرجع	تعداد سند مشکوک
۱۵۰۰	۳۰۰

ارزیابی – مجموعه داده فارسی (ترازبندی متن)

تعداد جفت سند	سطح سرقت علمی
۶۵۲	بدون سرقت علمی
۱۲۹	سرقت علمی بدون ابهام
۲۸۲	سرقت علمی با ابهام تصادفی

ارزیابی – مجموعه داده انگلیسی (ترازبندی متن)

★ مجموعه داده آموزش و آزمایش

تعداد جفت سند	سطح سرقت علمی
۱۰۰۰	بدون سرقت علمی
۱۰۰۰	سرقت علمی بدون ابهام
۱۰۰۰	سرقت علمی با ابهام تصادفی
۱۰۰۰	سرقت علمی با ابهام ترجمه
۱۱۸۵	سرقت علمی با ابهام خلاصه

$$\text{دقت} = \frac{\text{تعداد سندهای درست بازیابی شده}}{\text{تعداد سندهای اشتباه بازیابی شده} + \text{تعداد سندهای درست بازیابی شده}} \quad (\star)$$

$$\text{فراخوانی} = \frac{\text{تعداد سندهای بازیابی شده}}{\text{تعداد سندهای بازیابی نشده} + \text{تعداد سندهای بازیابی شده}}$$

$$F \text{ معیار} = 2 * \frac{\text{فراخوانی} * \text{دقت}}{\text{فراخوانی} + \text{دقت}}$$

معیارهای ارزیابی - ترازبندی متن

$$\text{دقت} = \frac{\text{تعداد بخش‌های درست بازیابی شده}}{\text{تعداد بخش‌های درست بازیابی شده} + \text{تعداد بخش‌های اشتباه بازیابی شده}}$$

$$\text{فراخوانی} = \frac{\text{تعداد بخش‌های بازیابی شده}}{\text{تعداد بخش‌های بازیابی نشده} + \text{تعداد بخش‌های بازیابی شده}}$$

$$F \text{ معیار} = 2 * \frac{\text{فراخوانی} * \text{دقت}}{\text{فراخوانی} + \text{دقت}}$$

کد ارزیابی در سایت مسابقات PAN
قرار داده شده است.

⊛ نتایج بدست آمده برای تشخیص سرقت علمی اسناد فارسی

⊛ ۱- تاثیر پارامتر Ω

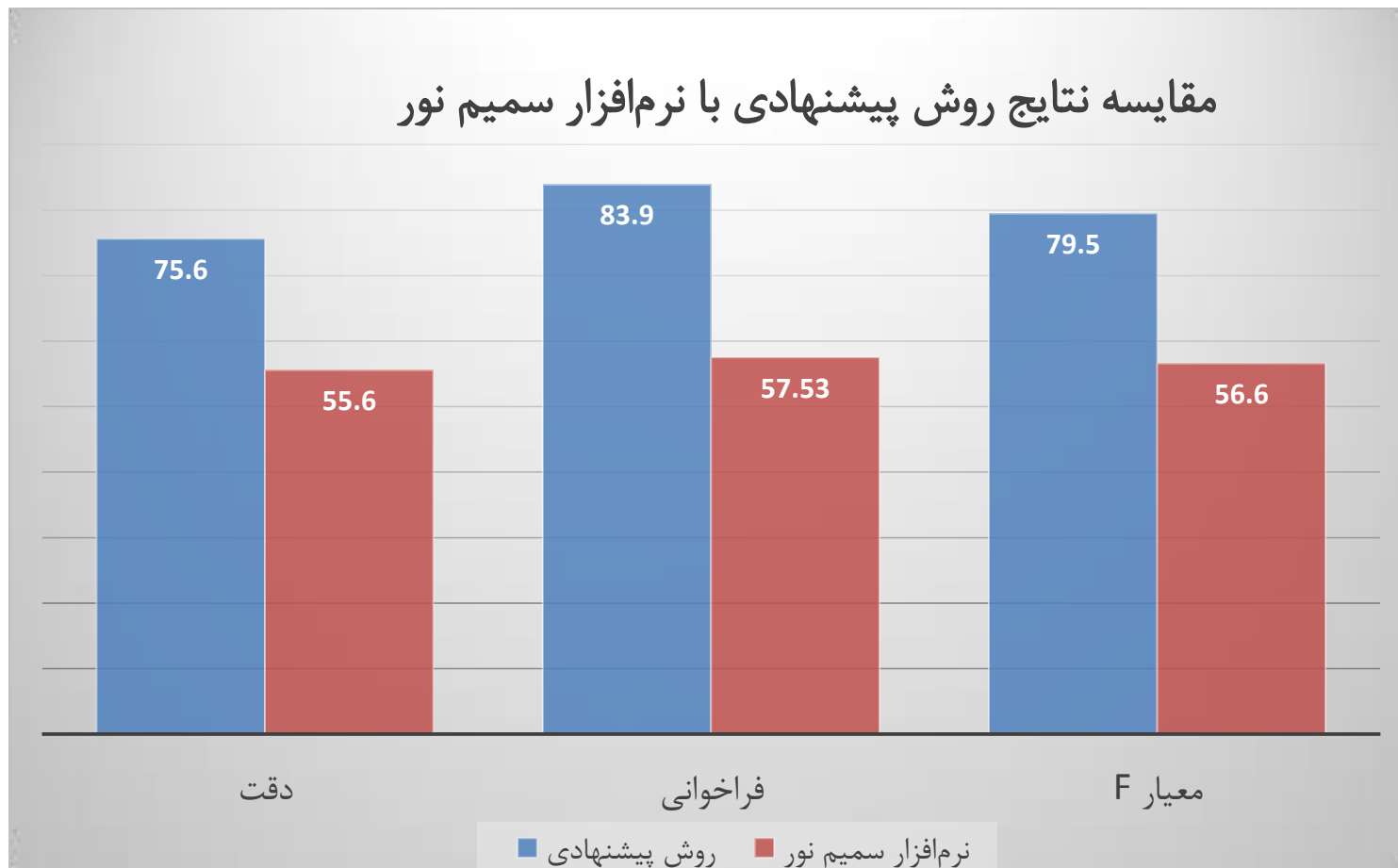
⊛ ۲- سطح بازیابی منبع

⊛ ۳- سطح ترازبندی متن

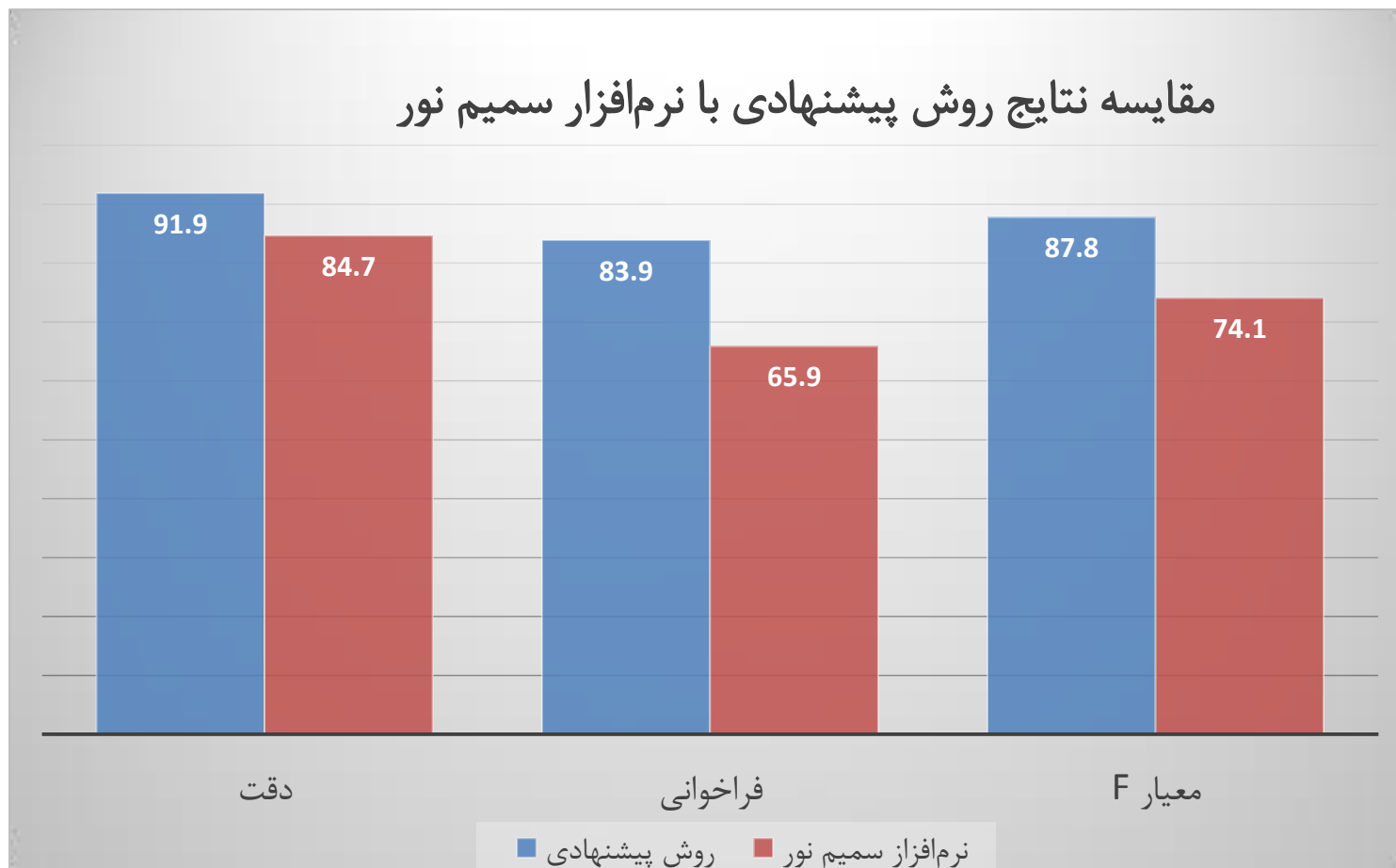
(تعداد دوتایی‌های مشترک) متفاوت در مجموعه Ω کارایی روش پیشنهادی براساس مقادیر مختلف داده سرقت معنایی و با درجه ابهام بالا

معیار F	فراخوانی	دقت	Ω حد آستانه
۶۳.۱۷	۴۸.۷۵	۸۹.۷۰	$\Omega = 20$
۶۸.۴۸	۵۵.۶۷	۸۸.۹۵	$\Omega = 15$
۷۴.۸۷	۶۶.۹۶	۸۴.۹۱	$\Omega = 10$
۷۹.۵	۸۳.۹	۷۵.۶	$\Omega = 6$

⊛ مجموعه داده سرقت معنایی و با درجه ابهام بالا



⊛ مجموعه داده همه سطوح سرقت علمی





ارزیابی – نتایج بخش ترازبندی سند اسناد فارسی

معیار F	فراخوانی	دقت	
۹۲.۵۷	۹۵.۷۱	۹۴.۰۰	سرقت علمی بدون ابهام
۸۹.۸۴	۸۶.۲۶	۹۳.۷۳	سرقت علمی با ابهام
۹۰.۰۰	۸۹.۶۶	۹۰.۷۴	کل پیکره

- ⊛ نتایج بدست آمده برای تشخیص سرقت علمی اسناد انگلیسی
- ⊛ ۱- بررسی تاثیر چند پارامتر مهم در تشخیص سرقت علمی نوع خلاصه
- ⊛ ۲- همه سطوح ترازبندی متن



تأثیر پارامترهای مختلف در تشخیص سرقت علمی با ابهام خلاصه

⊛ تاثیر جهت‌دار بودن گراف متن در تشخیص سرقت علمی نوع خلاصه اسناد انگلیسی

نوع گراف	دقت	فراخوانی	معیار F
گراف جهت‌دار	۸۷.۴۴	۸۰.۵۹	۸۳.۸۸
گراف غیر جهت‌دار	۷۸.۷۶	۸۱.۴۰	۸۰.۳۶

تأثیر پارامترهای مختلف در تشخیص سرقت علمی با ابهام خلاصه-۱

⊛ تاثیر حداقل آستانه مشابهت گره‌ها در تشخیص سرقت علمی نوع خلاصه اسناد انگلیسی

نوع گراف	دقت	فراخوانی	معیار F
$\alpha=0.4$	۶۱.۷۷	۸۴.۴۴	۷۱.۳۴
$\alpha=0.6$	۸۷.۴۴	۸۰.۵۹	۸۳.۸۸
$\alpha=0.8$	۸۷.۷۵	۶۰.۹۳	۶۴.۴۵



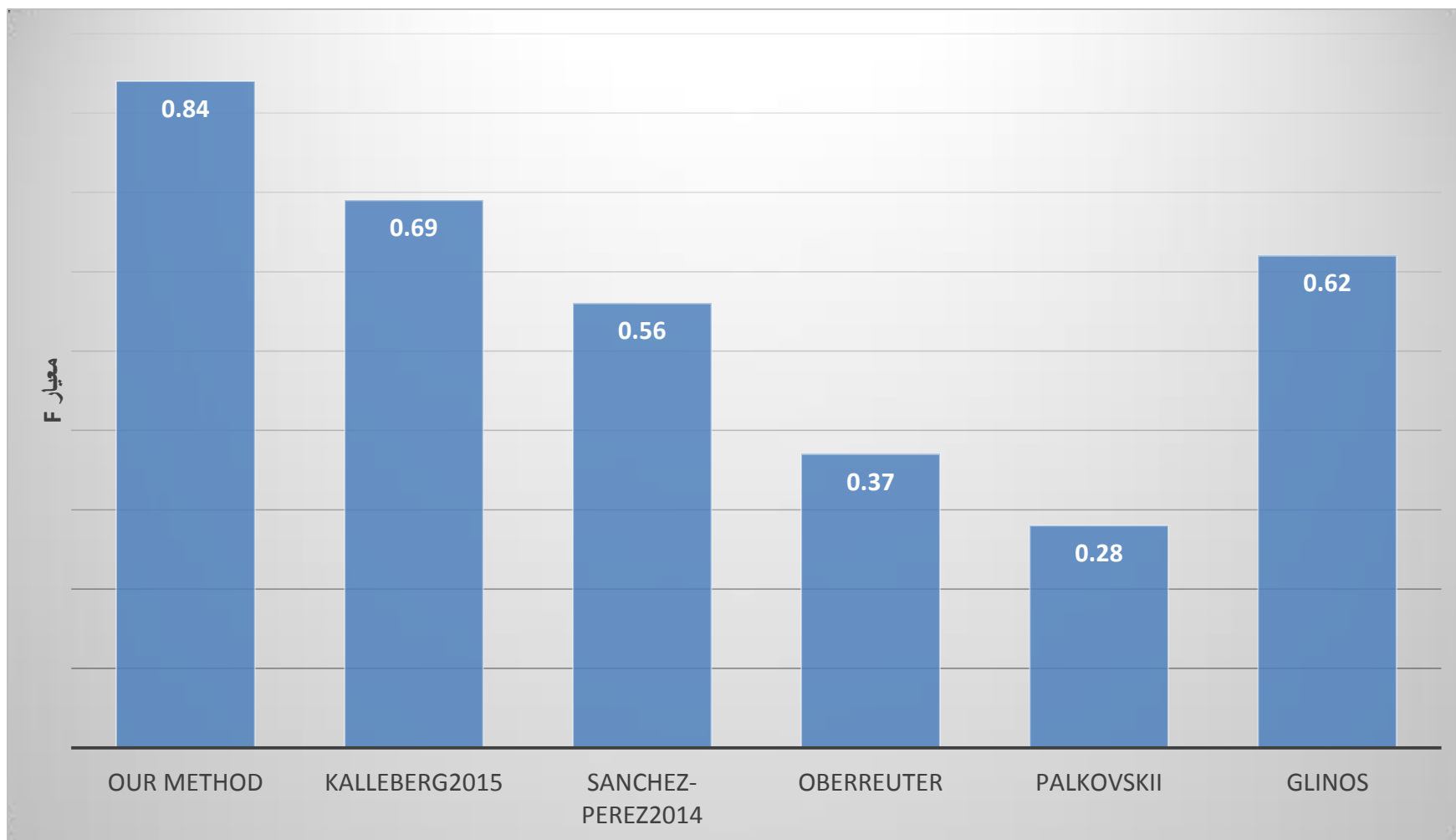
تأثیر پارامترهای مختلف در تشخیص سرقت علمی با ابهام خلاصه-۲

⊛ تأثیر نوع جداسازی متن در تشخیص سرقت علمی نوع خلاصه اسناد انگلیسی

نوع گراف	دقت	فراخوانی	معیار F
تبدیل متن به مجموعه‌ای از پاراگراف‌ها	۸۷.۴۴	۸۰.۵۹	۸۳.۸۸
تبدیل متن به مجموعه‌ای از جملات	۹۱.۷۱	۱۵.۱۱	۲۷.۱۷
تبدیل متن به بخش‌هایی با طول ۱۰۰۰ کاراکتر	۶۶.۵۶	۱۴.۳۱	۲۳.۵۶
تبدیل متن به بخش‌هایی با طول ۵۰۰ کاراکتر	۷۶.۹۴	۸.۸۱	۱۵.۸۲
تبدیل متن به بخش‌هایی با طول ۳۰۰ کاراکتر	۸۳.۳۷	۷.۵۰	۱۳.۹۰

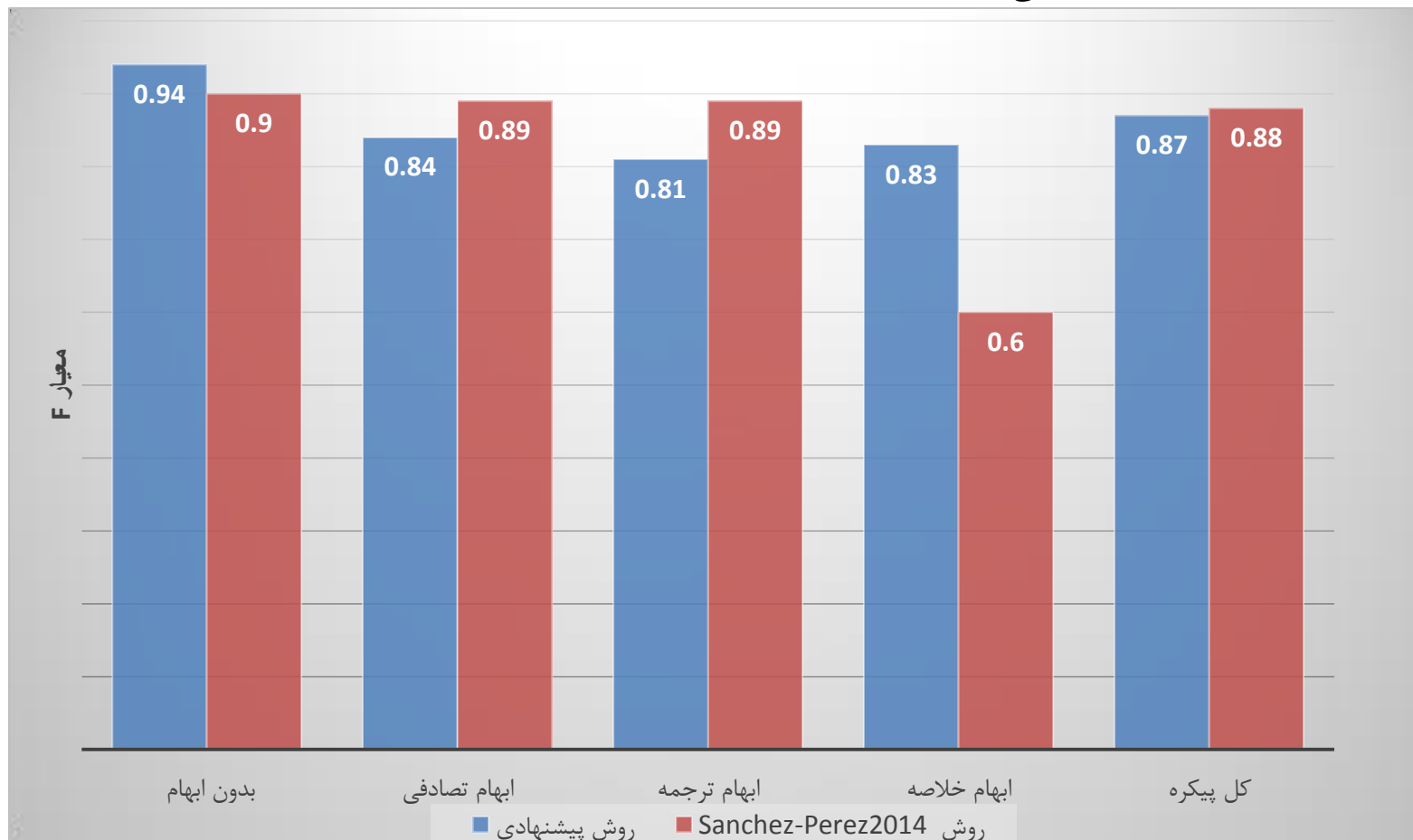


مقایسه عملکرد روش پیشنهادی در تشخیص سرقت علمی با ابهام خلاصه با روش‌های پیشین



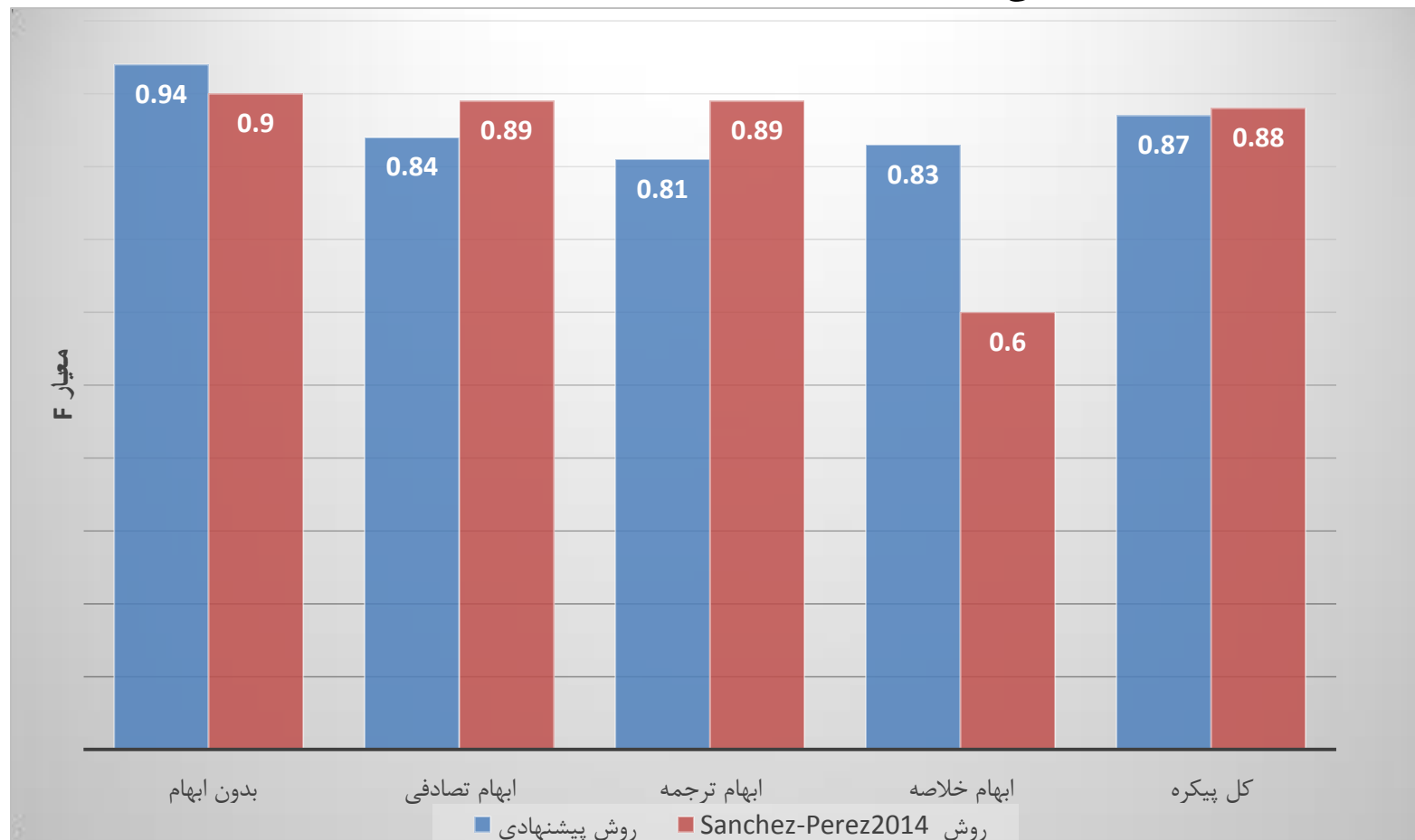
ارزیابی - نتایج بخش ترازبندی سند اسناد انگلیسی

⊙ مجموعه داده آموزش



ارزیابی - نتایج بخش ترازبندی سند اسناد انگلیسی

⊙ مجموعه داده آزمایش





⊛ مجموعه داده آموزش:

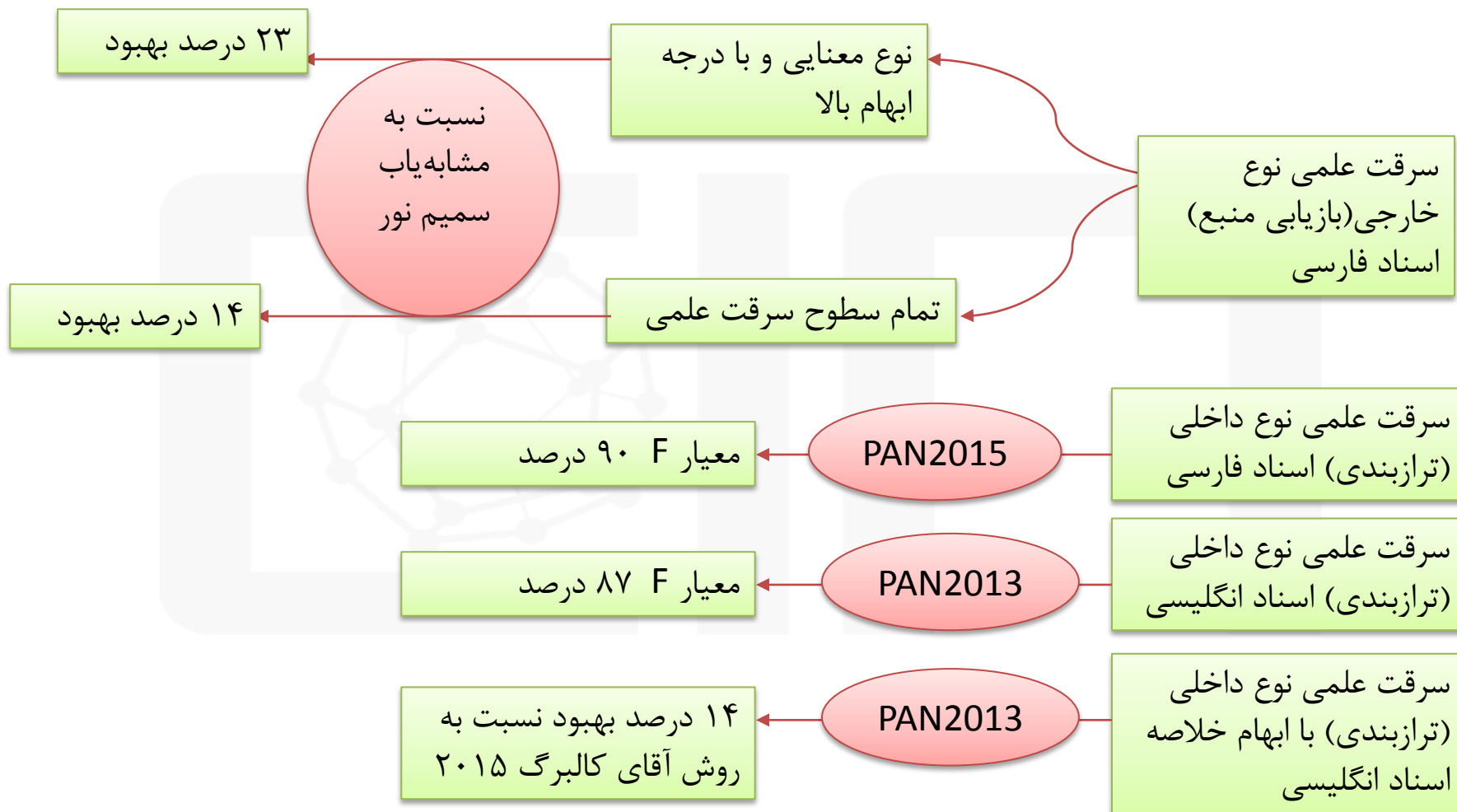
جدول ۲- نتایج به‌دست آمده از اجرای روش پیشنهادی روی مجموعه داده PAN2016

معیار F	فراخوانی	دقت	
۸۹.۳	۸۸.۵	۹۰	نتایج



⊙ مجموعه داده آزمایش:

Rank	Team	Plagdet	Granularity	Precision	Recall
1	Fatemeh Mashhadi, Mehrnoush Shamsfard Shahid Beheshti University, NLP Research Lab	0.92204	1.00146	0.92688	0.91919
2	Hadi Veisi, Kayvan Bijari, Kiarash Zahirnia, Erfaneh Gharavi University of Tehran, Data & Signal processing Lab	0.90593	1	0.95927	0.85820
3	Mozhgan Momtaz, Kayvan Bijari, Davood Heidarpour University of Tehran, COIN Lab	0.87103	1	0.89258	0.85049
4	Mahdi Niknam University of Qom	0.83015	1.03968	0.92034	0.79602
5	Faezeh Esteki, Faramarz Safi Esfahani Najafabad Branch, Islamic Azad University	0.80083	1.0	0.93337	0.70124



⊛ بهبود دقت سرقت علمی معنایی با نسخه توسعه یافته شبکه واژگان فارس‌نت

⊛ انتخاب گره‌های موثرتر به کمک پیکره اسامی خاص

⊛ بهبود دقت به کمک از بین بردن اثر کلمات مرکب پرتکرار زبان فارسی

⊛ تعمیم روش به تشخیص سرقت علمی چند زبانه

- [6] Zini, Manuel, Fabbri, Marco, Moneglia, Massimo, and Panunzi, Alessandro, Plagiarism detection through multilevel text comparison, Second International Conference, IEEE, 2006.
- [7] Potthast, Martin, Gollub, Tim, Hagen, Matthias, Tippmann, Martin, Kiesel, Johannes, Rosso, Paolo, Stamatatos, Efstathios, and Stein, Benno. Overview of the 5th International Competition on Plagiarism Detection. in Forner, Pamela, Navigli, Roberto, and Tufis, Dan, eds. , Working Notes Papers of the CLEF 2013 Evaluation Labs, September 2013.
- [8] Zager, Laura. Graph similarity and matching. Ph.D. thesis, Massachusetts Institute of Technology, 2005.
- [9] Blanco, Roi and Lioma, Christina. Graph-based term weighting for information retrieval. Information retrieval, 15(1):54–92, 2012.
- [10] Osman, Ahmed Hamza, Salim, Naomie, and Binwahlan, Mohammed Salem. Plagiarism detection using graph-based representation. arXiv preprint arXiv:1004.4449, 2010.
- [11] K. e. a. Khoshnavataher, "Developing Monolingual Persian Corpus for Extrinsic Plagiarism Detection Using Artificial Obfuscation," PAN-CELF, 2015.

- [1] N.Kumar, A graph based automatic plagiarism detection technique to handle artificial word reordering and paraphrasing, Computational Linguistics and Intelligent Text Processing, p. 481–494, 2014.
- [2] Sonawane, S. S., & Kulkarni, P. A, "Graph based Representation and Analysis of Text Document: A Survey of Techniques.," in International Journal of Computer Applications, 96(19), 2014.
- [3] Kalleberg, Rune Borge. Towards detecting textual plagiarism using machine learning methods. 2015.
- [4] Sanchez-Perez, Miguel A, Sidorov, Grigori, and Gelbukh, Alexander F. A winning approach to text alignment for text reuse detection at pan 2014. in CLEF (Working Notes), pp. 1004–1011, 2014.
- [5] Oberreuter, Gabriel, and Juan D. Velásquez, Text mining applied to plagiarism detection: The use of words for detecting deviations in the writing style, Expert Systems with Applications 40.9: 3756-3763, 2013.

از توجه شما عزیزان سپاسگزارم

سرانجام روزی به حکمت
همه اتفاقات زندگیتان
پی خواهید برد. پس
فعلاً به سردرگمی‌ها
بخندید، از میان اشکها
لیخند بزنید، و همواره
به خودتان یادآوری
کنید که پشت هر
حادثه‌ای دلیلی نهفته
است.