



بعد فراکتالی واژگان در یک نوشته

امیرحسین درونه *

پیشگفتار

پیشرفت‌های فناوری در دو دهه گذشته سبب دسترسی ما به انبوهی از اطلاعات شده است که بیشتر آنها اسنادی نوشتاری هستند. یافتن اطلاعات مورد نیاز به صورت دستی در این انبوه، کاری دشوار است. روش‌های گوناگونی تاکنون معرفی شده‌اند تا بتوان اطلاعات مهم یک نوشته یا شماری از نوشته‌ها را به شکل خودکار بیرون کشید. مجموعه این روش‌ها نوشته کاوی^۱ نامیده می‌شود. نوشته کاوی موضوعی میان رشته‌ای است که نخست در علوم رایانه مطرح شد و سپس دانش‌پیشگان دیگر حیطه‌ها مانند زبان‌شناسی محاسباتی، ریاضی، آمار و فیزیک نیز به آن علاقه‌مند شدند. رتبه بندی واژگان بر پایه اهمیت آنها در یک نوشته از موضوعات مهم نوشته کاوی است که می‌توان از آن در یافتن کلید واژگان، چکیده‌سازی نوشته، دسته‌بندی نوشته‌ها و دیگر موضوعات نوشته کاوی سود برد. تاکنون روش‌های گوناگونی برای رتبه‌بندی واژگان پیشنهاد شده‌اند.

اگر واژه‌های یک نوشته را بر حسب فراوانی (تعداد تکرار) آنها مرتب کنیم. به فراوان‌ترین واژه رتبه یک و دومین واژه فراوان رتبه دو و به همین ترتیب به باقی واژه‌ها رتبه اختصاص دهیم. در این فهرست مرتب شده از واژگان، واژه‌هایی مانند *of, the* و *a* برای زبان انگلیسی در بالای فهرست قرار می‌گیرند. این واژه‌ها در تمام نوشته‌ها یافت می‌شوند و مفهوم خاصی ندارند. اما واژه‌های پر اهمیت که به موضوع نوشته ربط دارند و معنای مورد نظر نویسنده را به خواننده منتقل می‌کنند معمولاً در رتبه‌های میانی قرار می‌گیرند. می‌توان از این واقعیت برای استخراج واژه‌های پر اهمیت سود برد [۱]. وجود واژه‌های کم اهمیت در رتبه‌های میانی دقت این روش را کاهش می‌دهد. افزون بر این محدوده میانی فهرست تعریف دقیقی ندارد. به صورت تجربی دیده شده است که واژه‌های مهم نوشته خوشه تشکیل می‌دهند یعنی در برخی از جاهای نوشته نزدیک یکدیگر دیده می‌شوند و توزیع یکنواختی در طول نوشته ندارند [۲]. اگر فاصله بین ظهور

متوالی یک واژه را بر حسب تعداد واژه میان آنها در نوشته اندازه‌گیری کنیم، می‌توان میزان افت و خیز آماری فواصل ظاهر شدن یک واژه را همچون سنج‌های برای اندازه‌گیری میزان خوشه شدن آن و در نتیجه رتبه‌بندی واژگان بکار برد [۳]. با استفاده از مکانیک آماری می‌توان نشان داد تابع توزیع فواصل ظهور متوالی یک واژه شکل $(1 + (1 - q)\beta x)^{\frac{1}{1-q}}$ دارد که در آن x متغیر فاصله، q و β دو پارامتر حقیقی هستند. می‌توان پارامتر q را معیاری برای اندازه‌گیری خوشه شدن دانست و واژگان یک نوشته را بر حسب آن رتبه‌بندی کرد [۴]. ویژگی خوشه‌بندی سبب می‌شود فراوانی واژه در بخش‌های گوناگون نوشته یکسان نباشد. آنتروپی سنج خوبی برای اندازه‌گیری این ناهمسانی است. آنتروپی به شکل $-\sum_i f_i(w) \ln f_i(w)$ تعریف می‌شود که در آن $f_i(w)$ نسبت فراوانی واژه w در بخش i -ام به فراوانی همان واژه در کل نوشته است. از آنتروپی برای جداسازی نوشته‌های طبیعی و کاتوره‌ای (تصادفی) استفاده شده است [۵، ۶]. گونه‌های دیگر آنتروپی نیز که بر حسب توزیع فواصل واژگان هستند نیز برای رتبه‌بندی پیشنهاد شده است [۷]. هر نوشته را می‌توان با یک گراف نمایش داد. مجموعه واژه‌های مستقل رأس‌های گراف هستند. مجاورت دو واژه در نوشته را با یالی میان رأس‌های آنان در گراف نشان می‌دهیم. با کاربست الگوریتم رتبه پیچ^۲ بر روی گراف وابسته به یک نوشته، می‌توان واژگان آن را رتبه‌بندی کرد [۸]. این الگوریتم همانی است که گوگل از آن برای مرتب‌سازی نتایج جستجوی خود سود می‌برد.

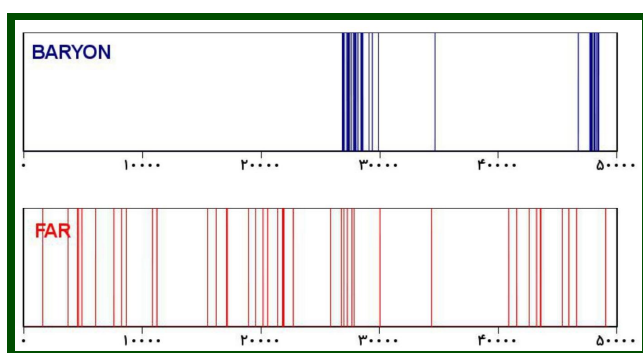
در این نوشتار ما با دیدگاهی نو به رتبه‌بندی واژگان می‌پردازیم. بر پایه این دیدگاه فرض می‌کنیم مکان‌هایی که یک واژه در نوشته ظاهر می‌شود الگویی خودهمانند را پدید می‌آورد. هر الگوی خود همانند دارای بعدی فراکتالی است که می‌توان به صورت عملی آن را محاسبه کرد. در یک نوشته که به گونه‌ای تصادفی درهم‌ریخته^۳ شود بعد فراکتالی همه واژه‌ها نزدیک یک است اما در نوشته‌ای که دارای

^۱Text Mining ^۲Page Rank ^۳Shuffled Text

حسب s رسم کنیم. شیب نمودار، اندازه D را بدست می‌دهد. در بخش بعدی درستی این فرض را برای واژگان یک کتاب نمونه می‌آزماییم.

positrons was	almost exactly	equal to	the number	of electrons	in addition	to electrons	and positrons
positrons was almost exactly	equal to the number	of electrons in addition	to electrons and positrons				
positrons was almost exactly equal to the number	of electrons in addition to electrons and positrons						

شکل ۱: سیمایی از جعبه‌بندی برای بخشی از نوشته. در بالا اندازه جعبه ۲ و در لایه میانی و پایینی به ترتیب ۴ و ۸ می‌باشد.



شکل ۲: نمای طیفی دو واژه BARYON و FAR. خوشه شدن و توزیع یکنواخت را برای این دو واژه می‌توان دید. هر دو واژه ۴۴ بار در نوشته تکرار شده‌اند. خوشه شدن و توزیع یکنواخت را برای این دو واژه می‌توان دید. هر دو واژه ۴۴ بار در نوشته تکرار شده‌اند.

اندازه‌گیری بعد فراکتالی واژگان

در این بخش کتاب ارزشمند؛ سه دقیقه نخست؛ نوشته استیون واینبرگ [۹] را برای بررسی برمی‌گزینیم. این کتاب درباره آغاز کیهان بر پایه نظریه مه‌بانگ (انفجار بزرگ) ^۶ و الگوی متعارف ^۷ حاکم بر برهم‌کنش ذرات بنیادی است و چگونگی پیدایش ماده را در سه دقیقه آغاز جهان شرح می‌دهد. طول این کتاب ۴۹۸۶ واژه است و واژه‌نامه‌ای با اندازه ۴۰۳۹ واژه دارد. اندازه جعبه‌ها را توانی از ۲ برمی‌گزینیم. شکل سیمایی از جعبه‌ها را برای بخشی از این نوشته نشان می‌دهد.

چیدمان منحصر به فردی از واژگان در نوشته است که معنای دلخواه نویسنده را به خواننده منتقل می‌کند. هرگونه به هم ریختگی در این چینش معنای آن را از بین می‌برد. اگر توزیع مکانی واژه‌ای با به هم ریختن چیدمان واژگان تغییر چشمگیری نداشته باشد بدین معنی است که آن واژه اهمیت کمی در انتقال معنای نوشته دارد. در نوشته‌ای که به صورت تصادفی به هم ریخته شود، واژگان همه

مفهوم است، واژه‌هایی که بار معنایی نوشته را به دوش می‌کشند دارای بعد فراکتالی کمتر از یک هستند. بعد فراکتالی واژه می‌تواند همچون سنجه‌ای برای رتبه‌بندی واژگان بکار آید.

در بخش بعدی در این باره بیشتر خواهیم گفت و در بخش پایانی از داده‌های تجربی برای پشتیبانی از این فرض استفاده می‌کنیم.

نوشته‌ها و فراکتال‌ها

زبان بشری سامانه‌ای پیچیده است که واژه‌ها اجزای بنیادین آن هستند. زبان سامانه‌ای بسته نیست بلکه ارتباط تنگاتنگی با پیرامون خود دارد. با گذشت زمان و ساخته شدن مفاهیم جدید تعداد واژگان یک زبان تغییر می‌کند و در گذر زمان حتی معنی برخی از واژه‌ها نیز دگرگون می‌شود. نوشته‌ها نیز بخشی مهم از زبان هستند که پیچیدگی آن را به ارث می‌برند. توزیع فراوانی واژه‌ها، رشد تعداد واژه‌های مستقل با افزایش طول نوشته و برخی دیگر از ویژگی‌های آماری، نمونه‌های این پیچیدگی هستند. در یک نوشته واژه‌ها در حال برهم‌کنش با یکدیگرند. این برهم‌کنش به دو گونه دستوری و معنایی است. برهم‌کنش دستوری در یک نوشته سبب می‌شود که واژه‌ها در یک جمله به ترتیب ویژه‌ای کنار یکدیگر قرار گیرند. این برهم‌کنش کوتاه برد است و برد آن اندازه یک یا دو جمله است. برهم‌کنش بین برخی از واژه‌ها در نوشته برای انتقال مفهوم به خواننده، برهم‌کنش معنایی نام دارد. این برهم‌کنش بلند برد است و چینش واژه‌ها را در سراسر نوشته هماهنگ می‌کند. برهم‌کنش معنایی سبب پیدایش الگویی فضایی برای هر واژه می‌شود. در این نوشتار فرض می‌کنیم این الگو دارای ویژگی خود همانندی ^۴ یا فراکتالی است، همان گونه که در بسیاری از سامانه‌های پیچیده دیده می‌شود. بعد این فراکتال‌ها را می‌توان با روش جعبه‌شماری ^۵ اندازه گرفت. برای این کار نوشته را به صورت دنباله یک بعدی از واژگان می‌انگاریم. هر s واژه متوالی را درون یک جعبه فرضی قرار می‌دهیم. بدین گونه نوشته به N_s جعبه تقسیم می‌شود. شکل نمایی از این جعبه‌بندی را نشان می‌دهد. شمار جعبه‌هایی را که یک واژه معین در آن یافت می‌شود را با نماد $N_w(s)$ نشان می‌دهیم. اگر توزیع هر واژه الگویی خود همانند داشته باشد، برای اندازه‌های بزرگ s خواهیم داشت:

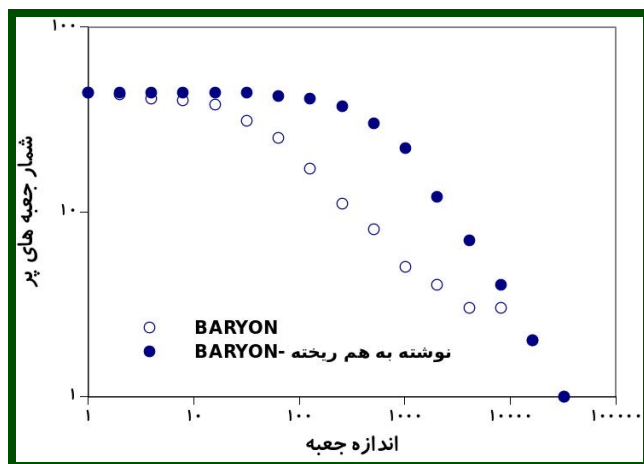
$$N_w(s) \approx s^{-D} \quad (1.2)$$

که D بعد فراکتالی واژه نام دارد.

برای محاسبه D کافی است نمودار تمام لگاریتمی $N_w(s)$ را بر

⁴S ⁵Box Counting ⁶Big Bang ⁷Standard Model

از واژه BARYON بیشتر است. برازش انجام شده نشان می‌دهد بعد واژه نخست نزدیک ۱ و دومی برابر ۰/۶ است.

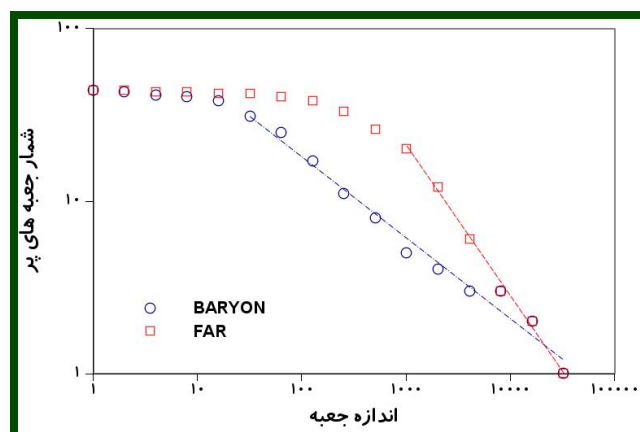


شکل ۴: BARYON در نوشته اصلی و به هم ریخته. همان گونه که دیده می‌شود در نوشته به هم ریخته بعد فراکتالی این واژه نزدیک ۱ است و با بعد فراکتالی آن در نوشته اصلی تفاوت چشمگیری دارد.

شکل مقایسه‌ای است بین شمار جعبه‌های پر برای واژه پر اهمیت BARYON در متن اصلی و متنی که به صورت تصادفی به هم ریخته شده است. در متن به هم ریخته بعد فراکتالی این واژه نزدیک ۱ است و در متن اصلی به صورتی چشمگیر کمتر از آن است.

همان گونه که در ابتدای این نوشته اشاره کردیم، تلاش ما برای یافتن راهی است که به صورتی خودکار واژگان را رتبه‌بندی کند. نوشتن الگوریتمی که خودکار بعد فراکتالی را محاسبه کند دشوار است. برای غلبه بر این دشواری به جای استفاده از بعد فراکتالی از کمیتی دیگر بهره می‌بریم که با آن رابطه دارد. این کمیت سطح بین نمودار تمام لگاریتمی جعبه‌شماری یک واژه در متن اصلی و به هم ریخته است. اگر واژه‌ای کم اهمیت باشد توزیع آن در نوشته اصلی و به هم ریخته تفاوت کمی با یکدیگر دارند که پیامد آن نزدیکی نمودارهای جعبه‌شماری آن در نوشته اصلی و به هم ریخته است. بدین گونه برای واژگان کم اهمیت اندازه این سطح ناچیز است. در شکل این دو سطح را برای دو واژه پر اهمیت BARYON و QUARKS نشان داده‌ایم.

توزیعی یکنواخت دارند بنابر این واژگان کم اهمیت نیز در نوشته اصلی دارای توزیعی یکنواخت هستند. شکل نمایی طیفی از دو واژه BARYON و FAR را نشان می‌دهد. هر دو واژه ۴۴ بار در نوشته تکرار شده‌اند. نخستین واژه در ارتباط با مفاهیم کتاب است و دومی واژه‌ای عمومی است که در هر نوشته‌ای می‌توان آن را یافت. جایگاه‌های ظاهر شدن این واژه‌ها در نوشته با خط‌های پررنگ مشخص شده‌اند. در این شکل، خوشه شدن برای واژه نخست و توزیع یکنواخت را برای دومین واژه می‌توان دید



شکل ۳: جعبه‌شماری برای واژگان BARYON و FAR. خط‌چین و خط‌چین نقطه برازش انجام شده برای بخش خطی نتایج جعبه‌شماری برای این دو واژه را در نوشته اصلی نشان می‌دهند.

فرض کنید در نوشته‌ای به طول N ، واژه کم اهمیت w دارای فراوانی M باشد. اگر این نوشته را با جعبه‌هایی به اندازه s واژه افزایش کنیم $N_s = N/s$ جعبه خواهیم داشت. برای s های کوچک ($N_s > M$) می‌دانیم $N_w(s) \simeq M$ است چون توزیع یکنواخت است و اثر خوشه‌بندی وجود ندارد که احتمال قرار گرفتن دو یا چند واژه در یک جعبه را افزایش دهد بنا بر این در هر جعبه یا یک واژه است و یا واژه‌ای یافت نمی‌شود. برای s های بزرگ ($N_s < M$) به ناچار همه جعبه‌ها دست کم یک رخداد از واژه را در خود باید جای دهند بدین گونه $N_w(s) = N_s$ می‌شود یعنی $N_w(s) \propto 1/s$ است. با توجه به معادله؟؟ بعد فراکتالی این دسته از واژگان برابر یک می‌شود. برای واژه‌های پر اهمیت اگر چه هنگامی که $s = 1$ است $N_w(s) = M$ می‌شود ولی با بزرگ شدن s به تندی از این مقدار فاصله می‌گیرد. برای s های بزرگ اثر خوشه‌بندی سبب می‌شود که $N_w(s) \ll N_s \propto 1/s$ باشد و بعد فراکتالی آن کمتر از یک شود. شکل پیامد جعبه‌شماری برای دو واژه BARYON و FAR را نشان می‌دهد. همان گونه که در شکل دیده می‌شود بعد فراکتالی واژه FAR

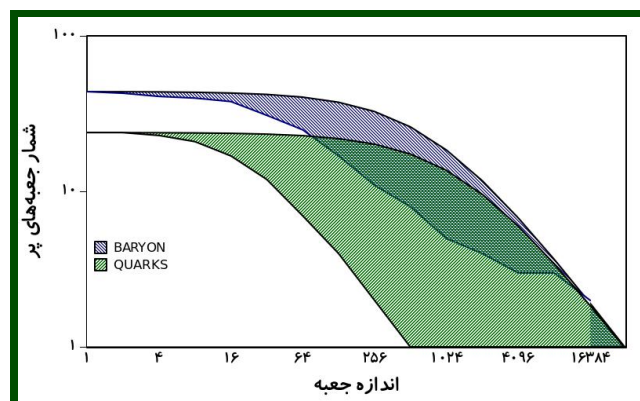
جدول ۱۰۲: فهرست ۲۰ واژه نخست که بر اساس اندازه سطحشان مرتب شده‌اند به همراه فراوانی آن‌ها در سمت راست و نیز فهرست مرتب شده از ۲۰ واژه نخست بر اساس فراوانی آن‌ها به همراه اندازه سطحشان در سمت چپ.

واژه	فراوانی	اندازه سطح	واژه	فراوانی	اندازه سطح
quarks	۲۴	۸۰۸۰.۱۷	the	۴۱۵۹	۳۹۷۴.۲
ions	۱۴	۸۴۸۸.۱۶	of	۲۴۴۷	۴۹۶۱.۲
diagram	۷	۲۹۵.۱۵	to	۱۱۵۴	۶۱۰۳۷.۲
rotating	۷	۷۲۷۲.۱۳	a	۱۰۷۸	۵۴۶۲.۲
quark	۱۰	۴۰۲۲.۱۳	in	۱۰۶۰	۵۵۷۷.۲
hydroxyl	۷	۸۷۹۹.۱۲	and	۹۴۵	۵۶۳۵.۲
nebulae	۱۴	۴۴۶۱.۱۲	is	۸۹۸	۴۹۲۶.۲
kmsec	۱۰	۲۳۹.۱۲	that	۷۳۱	۶۷۵۷.۲
antenna	۳۵	۲۲۷.۱۲	universe	۵۰۲	۲۱۵۷.۲
stone	۶	۱۴۹۳.۱۲	this	۴۳۴	۵۲۷۳.۲
slab	۴	۰۷۱۲.۱۲	it	۴۲۲	۵۵۶۳.۲
luminosity	۲۴	۵۰۸.۱۱	as	۴۱۸	۹۵۲۴.۲
cyanogen	۱۰	۳۸۷۳.۱۱	be	۳۹۵	۶۰۹۴.۲
angstroms	۴	۳۷۸.۱۱	at	۳۹۱	۵۳۲۳.۲
circumference	۴	۲۶۰۳.۱۱	by	۳۶۲	۸۱۷۳.۲
conservation	۲۳	۹۸۹۵.۱۰	we	۳۴۳	۹۰۴۹.۲
mc	۸	۴۷۴۵.۱۰	for	۳۲۵	۱۷۳۴.۲
planck	۱۸	۴۴۳۶.۱۰	are	۲۹۹	۸۶۲۳.۲
candidate	۲	۳۹۷۲.۱۰	was	۲۹۳	۱۲۷۵.۲
disc	۶	۳۷.۱۰	with	۲۸۹	۷۵۶۴.۲

علاقه‌مند را برای ارزیابی دقیق‌تر به مقاله نجفی و درونه [۱۰] ارجاع می‌دهیم.

مراجع

- [1] Luhn HP. The Automatic Creation of Literature Abstracts. IBM Journal of Research and Development, 2. 1958: 159–165.
doi: 10.1147/rd.22.0159
- [2] Ortuño M, Carpena P, Bernaola-Galvan P, Munoz E, So-moza AM. Keyword detection in natural languages and DNA. Europhysics Letters 57. 2002: 759–764.
doi:10.1209/epl/i2002-00528-3
- [3] Zhou H, Slater GW. A metric to search for relevant words. Physica A 329. 2003: 309–327.
doi: 10.1016/S0378-4371(03)00625-3
- [4] Mehri A, Daroonch AH. Keyword extraction by nonex-tensivity measure. Physical Review E 83. 2011:056106.



شکل ۵: سطح بین نمودارهای تمام لگاریتمی جعبه‌شماری برای دو واژه پر اهمیت BARYON و QUARKS. واژه دوم دارای فراوانی ۲۴ در نوشته است ولی اندازه سطح بزرگتری دارد.

با آن که واژه QUARKS تنها ۲۴ بار تکرار شده اما اندازه سطح بزرگتری دارد.

در جدول ۱۰۲ فهرستی از واژگان این کتاب دیده می‌شود که بر اساس اندازه سطحشان و نیز فهرستی از واژگانی که بر اساس فراوانی مرتب شده‌اند دیده می‌شوند. به سادگی می‌توان درستی این روش در رتبه‌بندی واژگان بر اساس اهمیت آن‌ها را تشخیص داد. خوانندگان

- Texts. Proceedings of conference on Empirical Methods in Natural Language Processing (EMNLP). 2004: 404–411.
- [9] Weinberg S. The First Three Minutes. Cambridge: Pegasus Press; 1993.
- [10] Najafi E, Darooneh AH. The Fractal Patterns of Words in a Text: A way for Automatic Keyword Extraction, PLoS ONE, 10(6). 2015:e0130617. doi:10.1371/journal.pone.0130617
- doi: 10.1103/PhysRevE.83.056106
- [5] Cohen A, Mantegna RN and Havlin S. Numerical analysis of word frequencies in artificial and natural language texts, Fractals 05. 1997: 95-104. doi:10.1142/S0218348X97000103
- [6] Herrera JP, Pury PA. Statistical keyword detection in literary corpora. Eur. Phys. J. B 63. 2008: 135. doi: 10.1140/epjb/e2008-00206-x
- [7] Mehri A, Darooneh AH. The role of entropy in word ranking. Physica A 390. 2011: 3157–3163. doi: 10.1016/j.physa.2011.04.013
- [8] Mihalcea R, Tarau P. TextRank: Bringing Order into
- * دانشگاه زنجان، گروه فیزیک

آگهی

ده سری پوستر رنگی: پنج سری به قطع 88×58 سانتی متر به نام های ابوریحان بیرونی، ابوالوفا بوزجانی، ابوعبدالله محمد بن موسی خوارزمی، غیاث الدین ابوالفتح عمر خیام و غیاث الدین جمشید کاشانی و پنج سری پوستر به قطع 68×48 سانتی متر به نام های تمدن اسلامی، دوران طلایی یونان، دوران های اولیه، عصر نوین و نوزائی (رنسانس)، از انتشارات ستاد ملی سال جهانی ریاضیات در دبیرخانه انجمن موجود است. بهای این ده پوستر ۱/۵۰۰/۰۰۰ ریال و هزینه ارسال آن ها ۳۰۰/۰۰۰ تعیین شده است.

این مجموعه زیبا و پرمحتوا می تواند زینت بخش کتابخانه ها، سالن ها، کلاس ها، اتاق ها و راهروهای دانشگاه ها، دبیرستان ها و مجامعی نظیر فرهنگ سراها و خانه های ریاضیات باشد.

از علاقه مندان، به ویژه مسئولان و مدیران محترم تقاضا می شود جهت خرید این مجموعه نفیس با دبیرخانه انجمن تماس بگیرند.