

MENINGKATKAN REKOMENDASI MENGGUNAKAN ALGORITMA PERBEDAAN TOPIK

Zen Munawar¹, Novianti Indah Putri², Dadad Zainal Musadad³

1. Manajemen Informatika, Politeknik LP3I Bandung
2. Teknik Informatika, Fakultas Teknologi Informasi Universitas Bale Bandung
3. Manajemen Informatika, Politeknik LP3I Bandung

ABSTRACT

This research addresses different topics, new methods designed to balance and disparate lists of personalized recommendations to reflect a complete spectrum of user interests. There is a flaw in the average accuracy, but the results suggest that this method can improve user satisfaction with recommended lists, particularly for lists created using an item-based collaborative filtering algorithm. This research was conducted based on previous research, namely the recommendation system, by taking into account the properties of the recommendation list as an entity, and specifically focusing on the accuracy of individual recommendations. Internal list similarity metrics were obtained to assess the diversity of recommended list topics and a topic difference approach to reduce internal list similarity. In this study, an evaluation method was carried out using book recommendation data, including offline analysis of rankings and online research.

Key Word: algorithms, recommendations, topics, measurements

ABSTRAK

Penelitian ini membahas perbedaan topik, metode baru yang dirancang untuk menyeimbangkan dan perbedaan daftar rekomendasi yang dipersonalisasi untuk mencerminkan spektrum minat pengguna yang lengkap. Terdapat kekurangan dalam keakuratan rata-rata, tetapi hasil menunjukkan bahwa metode ini dapat meningkatkan kepuasan pengguna dengan daftar rekomendasi, khususnya untuk daftar yang dibuat menggunakan algoritme pemfilteran kolaboratif berbasis item. Penelitian ini dilakukan berdasarkan pada penelitian terdahulu yaitu tentang sistem pemberi rekomendasi, dengan memperhatikan properti daftar rekomendasi sebagai entitas, juga secara khusus berfokus pada keakuratan rekomendasi individu. Diperoleh metrik kesamaan daftar internal untuk menilai keragaman topik daftar rekomendasi dan pendekatan perbedaan topik untuk mengurangi kesamaan daftar internal. Dalam penelitian ini dilakukan metode evaluasi dengan menggunakan data rekomendasi buku, termasuk analisis offline pada peringkat dan penelitian online.

Kata Kunci: algoritma, rekomendasi, topik, pengukuran

1. PENDAHULUAN

Sistem pemberi rekomendasi berguna untuk memberikan rekomendasi produk yang akan yang dipilih berdasarkan preferensi masa lalu, riwayat pembelian, dan informasi demografis. Banyak dari sistem yang paling sukses menggunakan pemfilteran kolaboratif [1] [2], dan banyak sistem komersial, misalnya, pemberi rekomendasi Amazon.com [3], memanfaatkan teknik ini untuk menawarkan daftar rekomendasi yang

dipersonalisasi kepada pelanggan. Secara komersial, e-commerce dapat disebut sebagai kegiatan yang berusaha menciptakan transaksi yang panjang antara perusahaan dan individu [4].

Meskipun akurasi sistem pemfilteran kolaboratif canggih, yaitu probabilitas bahwa pengguna aktif akan menghargai produk yang direkomendasikan sangat baik, beberapa implikasi yang mempengaruhi kepuasan pengguna telah diamati dalam implementasinya. Jadi, di Amazon.com, banyak rekomendasi yang

tampaknya "serupa" terkait konten. Misalnya, pelanggan yang telah membeli banyak prosa tertentu mungkin mendapatkan daftar rekomendasi di mana semua entri 5 teratas berisi buku yang ditulis oleh masing-masing penulisnya saja. Saat mempertimbangkan keakuratan, semua rekomendasi ini tampak sangat baik karena pengguna aktif sangat menghargai buku yang ditulis oleh penulis. Di sisi lain, dengan asumsi bahwa pengguna aktif memiliki beberapa minat selain penulis, misalnya, novel sejarah secara umum dan buku-buku tentang perjalanan dunia, kumpulan item yang direkomendasikan tampak buruk, karena kurangnya keragaman.

Secara tradisional, proyek sistem pemberi rekomendasi berfokus pada pengoptimalan akurasi menggunakan metrik seperti presisi / perolehan atau kesalahan absolut rata-rata. Sekarang penelitian telah mencapai titik di mana melampaui akurasi dan menuju pengalaman pengguna nyata menjadi sangat diperlukan untuk kemajuan lebih lanjut [5]. Rekomendasi bisa juga menimbulkan ketidakpastian yang disebut dengan certainty factor. Certainty factor adalah metode untuk mengelola ketidakpastian dalam sistem berbasis aturan [6].

Untuk mengatasi kekurangan yang disebutkan di atas dengan berfokus pada teknik yang berpusat pada kepuasan pengguna nyata daripada akurasi. Kontribusi dalam penelitian ini adalah sebagai berikut : Perbedaan topik, dengan mengusulkan pendekatan menyeimbangkan daftar rekomendasi N teratas sesuai dengan berbagai minat pengguna aktif. Metode ini mempertimbangkan keakuratan saran yang dibuat, dan tingkat minat pengguna pada topik tertentu. Analisis implikasi perbedaan topik pada pemfilteran kolaboratif berbasis pengguna [2] dan berbasis item [7] disediakan metrik kesamaan intra-daftar. Mengenai keragaman sebagai unsur penting untuk kepuasan pengguna, diperlukan metrik yang dapat mengukur fitur karakteristik tersebut. Metrik kesamaan intra-daftar sebagai cara yang efisien untuk

pengukuran, melengkapi metrik akurasi yang ada dalam upaya untuk menangkap kepuasan pengguna.

Akurasi versus kepuasan, Ada beberapa upaya di masa lalu yang menyatakan bahwa "akurasi tidak menceritakan keseluruhan cerita" [8]. Namun demikian, tidak ada bukti yang diberikan untuk menunjukkan bahwa beberapa aspek kepuasan pengguna yang sebenarnya melampaui akurasi. Analisis dari evaluasi online dan offline skala besar, mencocokkan hasil yang diperoleh dari metrik akurasi dengan kepuasan pengguna yang sebenarnya, dan menyelidiki interaksi dan penyimpangan antara kedua konsep.

2 KAJIAN TEORITIS

Pemfilteran Kolaboratif masih merupakan teknik yang paling umum diadopsi dalam menyusun sistem rekomendasi akademis dan komersial. Ide dasarnya mengacu pada membuat rekomendasi berdasarkan peringkat yang telah ditetapkan pengguna untuk produk. Saat ini, sejumlah besar data yang dikumpulkan dan dihasilkan setiap hari menawarkan berbagai peluang analitis bagi organisasi untuk mengungkap informasi yang bermanfaat untuk operasinya [9]. Pemeringkatan dapat bersifat eksplisit, yaitu dengan meminta pengguna menyatakan pendapatnya tentang produk tertentu, atau tersirat, ketika tindakan pembelian atau penyebutan suatu barang dihitung sebagai ekspresi apresiasi. Sementara peringkat implisit umumnya lebih mudah untuk dikumpulkan, penggunaannya menyiratkan menambahkan kebisingan ke informasi yang dikumpulkan [10].

2.1 Pemfilteran Kolaboratif Berbasis Pengguna

Pemfilteran Kolaboratif berbasis pengguna telah dieksplorasi secara mendalam selama sepuluh tahun terakhir [7] dan mewakili algoritme rekomendasi paling populer [2], karena kesederhanaannya yang menarik dan kualitas rekomendasi yang sangat baik. Pemfilteran Kolaboratif F beroperasi pada sekumpulan pengguna $A = \{a_1, a_2, \dots, a_n\}$,

satu set produk $B = \{b_1, b_2, \dots, b_m\}$, dan fungsi

pemeringkatan parsial $r_i: B! [-1, +1]^?$ untuk setiap pengguna a_i . Nilai negatif $r_i(b_k)$ menunjukkan ketidaksukaan, sedangkan nilai positif menunjukkan kesukaan a_i terhadap produk b_k . Jika peringkat hanya implisit, merepresentasikannya dengan himpunan $R_i B$, setara dengan $\{b_k \in B \mid r_i(b_k) \neq 0\}$.

Proses kerja Pemfilteran Kolaboratif berbasis pengguna dapat dibagi menjadi dua langkah utama yaitu pembentukan lingkungan. Dengan asumsi a_i sebagai pengguna aktif, nilai kesamaan $c(a_i, a_j) \in [-1, +1]$ untuk semua $a_j \in A \setminus \{a_i\}$ dihitung, berdasarkan kesamaan fungsi pemeringkatan masing-masing r_i, r_j . Secara umum, korelasi Pearson [11] atau jarak cosinus [2] digunakan untuk menghitung $c(a_i, a_j)$. M teratas pengguna yang paling mirip menjadi anggota lingkungan a_i .

Prediksi peringkat. Mengambil semua produk yang telah dinilai oleh a_i tetangga $a_j \in \text{klik}(a_i)$ dan yang baru untuk a_i , yaitu, $r_i(b_k) = ?$, Prediksi menyukai $w_i(b_k)$ dihasilkan. Nilai $w_i(b_k)$ dengan ini bergantung pada kesamaan $c(a_i, a_j)$ pemilih a_j dengan $r_j(b_k) \neq 0$, Serta peringkat $r_j(b_k)$ a_j tetangga ini ditetapkan ke b_k .

Akhirnya, daftar $Pw_i: \{1, 2, \dots, N\}!$ B dari rekomendasi atas- N dihitung, berdasarkan prediksi w_i . Perhatikan bahwa fungsi Pw_i bersifat injektif dan mencerminkan peringkat rekomendasi dalam urutan menurun, memberikan prediksi tertinggi terlebih dahulu.

2.2 Pemfilteran Kolaboratif Berbasis Item

Pemfilteran Kolaboratif berbasis item [7] telah mendapatkan momentum selama lima tahun terakhir berdasarkan karakteristik kompleksitas komputasi yang menguntungkan dan kemampuan untuk memisahkan proses komputasi model dari pembuatan prediksi aktual.

Khusus untuk kasus dimana $|A| \gg |B|$, kinerja komputasi Pemfilteran Kolaboratif berbasis item telah terbukti lebih unggul dari CF berbasis pengguna [7]. Keberhasilannya juga meluas ke banyak sistem pemberi rekomendasi komersial, seperti Amazon.com [3]. Seperti

Pemfilteran Kolaboratif berbasis pengguna, pembuatan rekomendasi didasarkan pada peringkat $r_i(b_k)$ bahwa pengguna $a_i \in A$ disediakan untuk produk $b_k \in B$. Namun, tidak seperti Pemfilteran Kolaboratif berbasis pengguna, nilai kesamaan c dihitung untuk item daripada pengguna, oleh karena itu $c: B \times B! [-1, +1]$. Secara kasar, dua item b_k, b_e adalah serupa, yaitu memiliki c besar (b_k, b_e), jika pengguna yang menilai salah satunya cenderung memberi peringkat yang lain, dan jika pengguna cenderung memberi mereka peringkat yang identik atau serupa. Selain itu, untuk setiap b_k , lingkungannya $(b_k) \in B$ dari M teratas item yang paling mirip ditentukan.

Prediksi $w_i(b_k)$ dihitung sebagai berikut:

$$w_i(b_k) = \frac{\sum_{b_e \in B_k^I} (c(b_k, b_e) \cdot r_i(b_e))}{\sum_{b_e \in B_k^I} |c(b_k, b_e)|} \quad (1)$$

Dimana

$$B_k^I := \{b_e \mid b_e \in \text{clique}(b_k) \wedge r_i(b_k) \neq 0\}$$

Secara intuitif, pendekatan tersebut mencoba meniru perilaku pengguna yang sebenarnya, meminta pengguna untuk menilai nilai dari produk yang tidak diketahui b_k dengan membandingkan yang terakhir dengan yang diketahui, item serupa dan mempertimbangkan seberapa besar a_i menghargainya. Penghitungan akhir dari daftar rekomendasi N teratas Pw_i mengikuti proses Pemfilteran Kolaboratif berbasis pengguna, menyusun rekomendasi menurut w_i dalam urutan menurun.

3 METRIK EVALUASI

Metrik evaluasi sangat penting untuk menilai kualitas dan kinerja sistem pemberi rekomendasi, meskipun masih dalam tahap awal. Sebagian besar evaluasi berkonsentrasi pada pengukuran akurasi saja dan mengabaikan faktor-faktor lain, misalnya, kebaruan dan kebetulan rekomendasi, dan keragaman item daftar yang direkomendasikan. Bagian berikut memberikan garis besar metrik populer. Sebuah survei ekstensif tentang metrik akurasi disediakan di [12].

3.1 Metrik Akurasi

Metrik akurasi telah ditentukan pertama dan terpenting untuk dua tugas utama:

Pertama, untuk menilai keakuratan prediksi tunggal, yaitu, seberapa banyak prediksi w_i (b_k) untuk produk yang menyimpang dari peringkat aktual a_i r_i (b_k). Metrik ini sangat cocok untuk tugas-tugas di mana prediksi ditampilkan bersama dengan produk, misalnya, penjelasan dalam konteks [12]. Kedua, metrik pendukung keputusan mengevaluasi keefektifan dalam membantu pengguna memilih item berkualitas tinggi dari rangkaian semua produk, umumnya mengangap preferensi biner.

Metrik akurasi prediksi mengukur seberapa dekat peringkat yang diprediksi menjadi peringkat pengguna sebenarnya. Paling menonjol dan banyak digunakan [2], *mean absolute error* (MAE) mewakili cara yang efisien untuk mengukur akurasi statistik dari prediksi w_i (b_k) untuk himpunan B_i produk:

$$|\overline{E}| = \frac{\sum_{b_k \in B_i} |r_i(b_k) - w_i(b_k)|}{|B_i|} \quad (2)$$

Terkait dengan MAE, *mean squared error* (MSE) mengkuadratkan kesalahan sebelum menjumlahkan. Karenanya, kesalahan besar menjadi jauh lebih jelas daripada kesalahan kecil. Sangat mudah diterapkan, metrik akurasi prediktif tidak sesuai untuk mengevaluasi kualitas daftar rekomendasi N teratas. Pengguna hanya peduli dengan kesalahan untuk produk peringkat tinggi. Di sisi lain, kesalahan prediksi untuk produk peringkat rendah tidak penting, karena pengguna tidak tertarik padanya. Namun, MAE dan MSE memperhitungkan kedua jenis kesalahan dengan cara yang persis sama.

Presisi dan perolehan, keduanya terkenal dari pengambilan informasi, tidak mempertimbangkan prediksi dan penyimpangannya dari peringkat sebenarnya. Umumnya lebih menyukai menilai seberapa relevan serangkaian rekomendasi yang diberi peringkat bagi pengguna aktif. Sebelum menggunakan metrik ini untuk validasi silang, K-folding diterapkan, membagi setiap produk yang diberi peringkat oleh pengguna b_k 2 R_i menjadi irisan K disjoint yang lebih disukai

berukuran sama. Dengan ini, K -1 irisan yang dipilih secara acak dari set pelatihan a_i R_{x_i} . Peringkat ini kemudian menentukan profil a_i dari mana rekomendasi akhir dihitung. Untuk pembuatan rekomendasi, potongan sisa a_i ($R_i \setminus R_{x_i}$) dipertahankan dan tidak digunakan untuk prediksi. Potongan ini, dilambangkan T_{x_i} , merupakan set pengujian, yaitu, produk yang ingin diprediksi oleh pemberi rekomendasi.

Sarwar [13] menyajikan varian penarikan yang disesuaikan, mencatat persentase produk set uji b 2 T_{x_i} yang terjadi dalam daftar rekomendasi P_{x_i} sehubungan dengan jumlah keseluruhan produk set uji | T_{x_i} |: dari produk set uji b 2 T_{x_i} yang terjadi dalam daftar rekomendasi P_{x_i} sehubungan dengan jumlah keseluruhan produk set uji | T_{x_i} |:

$$\text{Recall} = 100 \cdot \frac{|T_i^x \cap P_i^x|}{|T_i^x|} \quad (3)$$

Simbol = P_{x_i} adalah gambar peta P_{x_i} , yaitu semua item bagian dari daftar rekomendasi.

Dengan demikian, presisi mewakili persentase produk set uji b 2 T_{x_i} yang terjadi di P_{x_i} sehubungan dengan ukuran daftar rekomendasi:

$$\text{Precision} = 100 \cdot \frac{|T_i^x \cap P_i^x|}{|P_i^x|} \quad (4)$$

Breese dkk. [14] memperkenalkan ekstensi menarik untuk diingat, yang dikenal sebagai *recall* tertimbang atau skor Breese. Pendekatan ini memperhitungkan urutan daftar N teratas, menghukum rekomendasi yang tidak benar semakin ke bawah daftar itu terjadi. Penalti berkurang dengan peluruhan eksponensial.

Metrik pendukung keputusan populer lainnya termasuk ROC [15], "karakteristik operasi penerima". ROC mengukur sejauh mana sistem penyaringan informasi berhasil membedakan antara sinyal dan noise.

3.2 Melampaui Akurasi

Meskipun metrik akurasi merupakan aspek penting dari kegunaan, ada ciri-ciri

kepuasan pengguna yang tidak dapat mereka tangkap. Namun, metrik non-akurasi sebagian besar telah ditolak oleh minat penelitian utama sejauh ini.

Cakupan, Di antara semua metrik evaluasi non-akurasi, cakupan adalah yang paling sering digunakan [2]. Cakupan mengukur persentase elemen bagian dari domain masalah tempat prediksi dapat dibuat. Kebaruan dan Kebetulan. Beberapa pemberi rekomendasi memberikan hasil yang sangat akurat yang masih tidak berguna dalam praktiknya, misalnya, menyarankan pisang kepada pelanggan di toko bahan makanan. Meskipun sangat akurat, perhatikan bahwa hampir semua orang menyukai dan membeli pisang. Oleh karena itu, rekomendasi mereka tampak terlalu jelas dan tidak banyak membantu pembeli.

Metrik kebaruan dan kebetulan dengan demikian mengukur "ketidakjelasan" dari rekomendasi yang dibuat, menghindari "cherrypicking" [12]. Untuk ukuran kebetulan yang sederhana, ambillah popularitas rata-rata dari barang-barang yang direkomendasikan. Skor yang lebih rendah yang diperoleh menunjukkan kebetulan yang lebih tinggi.

3.3 Kesamaan IntraList

Metrik baru bertujuan untuk menangkap keragaman daftar. Dengan ini, keragaman dapat merujuk pada semua jenis fitur, misalnya genre, pengarang, dan karakteristik pembeda lainnya.

Berdasarkan fungsi kebebasan $c: B \times B! [-1, +1]$ mengukur kesamaan $c(b_k, b_e)$ antara produk b_k , sesuai dengan beberapa kriteria yang ditentukan khusus, mendefinisikan kesamaan intralis untuk daftar P_{wi} sebagai berikut:

$$ILS(P_{wi}) = \frac{\sum_{b_k \in \mathcal{P}_{P_{wi}}} \sum_{b_e \in \mathcal{P}_{P_{wi}}, b_k \neq b_e} c_o(b_k, b_e)}{2} \quad (5)$$

Skor yang lebih tinggi menunjukkan keragaman yang lebih rendah. Fitur matematika yang menarik dari $ILS(P_{wi})$ yang kita rujuk di bagian selanjutnya adalah permutasi-insensitivitas, yaitu, misalkan SN menjadi kelompok simetris

dari semua permutasi pada $N = |P_{wi}|$ simbol:

$$\forall \sigma_i, \sigma_j \in S_N: ILS(P_{wi} \circ \sigma_i) = ILS(P_{wi} \circ \sigma_j) \quad (6)$$

Oleh karena itu, hanya mengatur ulang posisi rekomendasi dalam daftar N teratas P_{wi} tidak mempengaruhi kesamaan intra-daftar P_{wi}

4 PERBEDAAN TOPIK

Salah satu masalah utama dengan metrik akurasi adalah ketidakmampuan mereka untuk menangkap aspek kepuasan pengguna yang lebih luas, menyembunyikan beberapa kekurangan yang mencolok dalam sistem yang ada [16]. Misalnya, menyarankan daftar item yang sangat mirip, misalnya, sehubungan dengan penulis, genre, atau topik, mungkin tidak banyak berguna bagi pengguna, meskipun akurasi rata-rata daftar ini mungkin tinggi.

Masalah ini telah dipahami oleh peneliti lain sebelumnya, yang disebut "efek portofolio" oleh Ali dan van Stam [17]. Sistem Pemfilteran Kolaboratif berdasarkan item secara khusus rentan terhadap efek tersebut. Laporan dari pemberi rekomendasi TV berbasis item TiVo [17], serta pengalaman pribadi dengan pemberi rekomendasi Amazon.com, juga berbasis item [3]. Misalnya, salah satu penulis ini hanya mendapat rekomendasi untuk buku Heinlein, yang lain mengeluh tentang semua buku yang disarankan adalah tulisan Tolkien.

Hukum menjelaskan efek kejenuhan yang terus menurunkan utilitas tambahan produk p ketika diperoleh atau dikonsumsi berulang kali. Misalnya, Anda ditawari minuman favorit Anda. Misalkan p_1 menunjukkan harga yang bersedia Anda bayarkan untuk produk itu. Dengan asumsi Anda ditawari segelas kedua minuman tersebut, jumlah p_2 uang yang ingin Anda belanjakan akan lebih rendah, yaitu $p_1 > p_2$. Sama untuk p_3, p_4 , dan lain sebagainya.

Pendekatan ini disebut perbedaan topik untuk menangani masalah yang dihadapi dan membuat daftar yang direkomendasikan lebih beragam dan dengan demikian lebih berguna. Metode ini

mewakili ekstensi untuk algoritma pemberi rekomendasi yang ada dan diterapkan di atas daftar rekomendasi.

4.1 Metrik Kesamaan Berbasis Taksonomi

Fungsi $c: 2B \times 2B! [-1, +1]$, mengukur kesamaan antara dua set produk, merupakan bagian penting dari perbedaan topik. Metrik menghitung kesamaan antara set produk berdasarkan klasifikasinya. Setiap produk termasuk dalam satu atau beberapa kelas yang disusun secara hierarkis dalam taksonomi klasifikasi, mendeskripsikan produk dengan cara yang dapat dibaca mesin.

Taksonomi klasifikasi ada untuk berbagai domain. Amazon.com membuat taksonomi yang sangat besar untuk buku, DVD, CD, barang elektronik, dan pakaian. Lihat Gambar 1 untuk satu contoh taksonomi. Selain itu, semua produk di Amazon.com memuat deskripsi konten yang berkaitan dengan taksonomi domain ini. Topik unggulan dapat mencakup penulis, genre, dan audiens.

4.2 Algoritma Perbedaan Topik

Algoritma 1 menunjukkan algoritma perbedaan topik lengkap, sketsa tekstual singkat diberikan di paragraf berikutnya.

Fungsi P_{w_i} menunjukkan daftar rekomendasi baru, hasil dari penerapan perbedaan topik. Untuk setiap entri daftar $z \in [2, N]$, mengumpulkan produk-produk b dari kandidat himpunan produk B_i yang tidak terjadi pada posisi $o < z$ di P_{w_i} dan menghitung kemiripannya dengan himpunan $\{P_{w_i}(k) \mid k \in [1, z]\}$, yang berisi semua rekomendasi baru sebelum peringkat z .

Menyortir semua produk b menurut $c(b)$ dalam urutan terbalik, kita memperoleh peringkat ketidaksamaan P_{rev}^c . Peringkat ini kemudian digabung dengan peringkat rekomendasi asli P_{w_i} menurut faktor perbedaan F , sehingga menghasilkan peringkat akhir P_{w_i} . Faktor F mendefinisikan dampak ketidaksamaan peringkat P_{rev}^c pada keluaran

keseluruhan akhirnya. Besar $F \in [0, 1]$ mendukung perbedaan atas urutan relevansi asli a_i , sedangkan $F \in [0, 0,5]$ menghasilkan daftar rekomendasi yang mendekati peringkat P_{w_i} asli. Untuk analisis eksperimental, menggunakan faktor perbedaan $F \in [0, 0,9]$. Perhatikan bahwa daftar masukan terurut P_{w_i} harus jauh lebih besar daripada daftar N teratas akhir. Untuk eksperimen selanjutnya, menggunakan daftar masukan 50 teratas untuk 10 rekomendasi teratas pada akhirnya

```

procedure diversify ( $P_{w_i}, \Theta_F$ ) {
     $B_i \leftarrow \mathcal{I}P_{w_i}; P_{w_i^*}(1) \leftarrow P_{w_i}(1);$ 
    for  $z \leftarrow 2$  to  $N$  do
        set  $B_i' \leftarrow B_i \setminus \{P_{w_i^*}(k) \mid k \in [1, z]\};$ 
         $\forall b \in B_i':$  compute  $c^*(b), \{P_{w_i^*}(k) \mid k \in [1, z]\}$ 
        compute  $P_{c^*}: \{1, 2, \dots, |B_i'|\}$ 
         $\rightarrow B_i'$  using  $C^*$ ;
        for all  $b \in B_i'$  do
             $P_{c^*}^{rev^{-1}}(b) \leftarrow |B_i'| - P_{c^*}^{-1}(b);$ 
             $w_i^*(b) \leftarrow P_{w_i}^{-1}(b) \cdot (1 - \Theta_F)$ 
        end do
         $P_{w_i^*}(z) \leftarrow \min\{w_i^*(b) \mid b \in B_i'\};$ 
    end do
    return  $P_{w_i^*};$ 
}

```

Algoritma 1: Perbedaan topik berurutan.

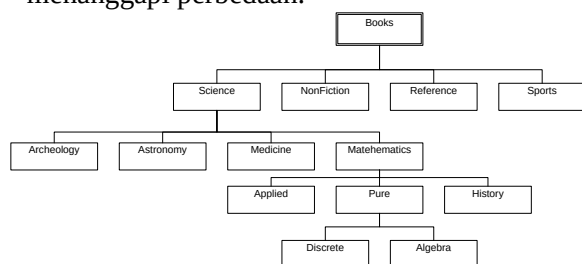
4.3 Ketergantungan Rekomendasi

Untuk bisa mengimplementasikan perbedaan topik, Diasumsikan bahwa produk yang direkomendasikan $P_{w_i}(o)$ dan $P_{w_i}(p)$, $o, p \in [2, N]$, bersama dengan deskripsi isinya, secara efektif memberikan dampak satu sama lain, yang biasanya diabaikan oleh pendekatan yang ada. : biasanya, hanya urutan bobot relevansi $o < p$ $w_i(P_{w_i}(o))$ $w_i(P_{w_i}(p))$ harus memegang untuk item daftar rekomendasi, tidak ada ketergantungan lain yang diasumsikan.

Dalam kasus perbedaan topik, saling ketergantungan rekomendasi berarti bahwa peringkat ketidaksamaan item b saat ini sehubungan dengan rekomendasi sebelumnya memainkan peran penting dan dapat memengaruhi peringkat baru.

5. ANALISIS EMPIRIS

Analisis dilakukan dengan cara evaluasi offline untuk memahami konsekuensi perbedaan topik pada metrik akurasi, dan analisis online untuk menyelidiki bagaimana metode ini memengaruhi kepuasan pengguna yang sebenarnya. Penerapan perbedaan topik dengan $F \in \{0, 0.1, 0.2, \dots, 0.9\}$ ke daftar yang dihasilkan oleh CF berbasis pengguna dan CF berbasis item, mengamati efek yang terjadi ketika terus meningkatkan F dan menganalisis bagaimana kedua pendekatan menanggapi perbedaan.



Gambar 1: Fragmen dari taksonomi buku Amazon.com

5.1 Perancangan Dataset

Perancangan dataset berdasarkan analisis online dan offline pada data yang dikumpulkan dari BookCrossing [18]. Dunia saat ini tidak lepas dari peran data karena semua dibangun di atas sebuah fondasi data [19]. Komunitas terakhir melayani pecinta buku yang bertukar buku di seluruh dunia dan berbagi pengalaman mereka dengan orang lain.

Pada tahap selanjutnya, dilakukan taksonomi buku Amazon.com, yang terdiri dari 13.525 topik berbeda. Untuk dapat menerapkan perbedaan topik, maka perlu menggali informasi konten, dengan fokus pada deskripsi taksonomi yang menghubungkan buku ke node taksonomi dari Amazon.com.

5.2 Eksperimen Offline

Eksperimen offline dilakukan dengan membandingkan presisi, perolehan, dan skor kesamaan intra-daftar untuk 20 penyajian daftar

rekomendasi yang berbeda. Separuh dari daftar rekomendasi ini didasarkan pada Pemfilteran Kolaboratif berbasis pengguna dengan derajat yang berbeda, yang lainnya pada Pemfilteran Kolaboratif berbasis item. Perhatikan bahwa penelitian ini tidak menghitung nilai metrik MAE karena dengan peringkat implisit daripada eksplisit.

Mesin pemberi rekomendasi berkinerja tinggi digunakan untuk menghitung daftar kasus dasar ini. Selanjutnya, menerapkan algoritme perbedaan ke kedua kasus dasar, menerapkan faktor F mulai dari 10% hingga 90%. Untuk evaluasi, semua daftar dipotong menjadi hanya 10 buku.

Hasil analisis, tertarik untuk melihat bagaimana akurasi, ditangkap oleh presisi dan recall, berperilaku ketika meningkatkan F . Karena perbedaan topik dapat membuat buku dengan akurasi yang diprediksi tinggi mengalir ke bawah daftar,

Precision dan *Recall*, Pertama, menganalisis presisi dan skor recall untuk kedua kasus dasar yang tidak terjadi perbedaan, yaitu, ketika $F = 0$. Tabel 1 menyatakan bahwa Pemfilteran Kolaboratif berbasis pengguna dan item-based menunjukkan akurasi yang hampir identik, yang ditunjukkan oleh nilai presisi. Nilai ingatan mereka sangat berbeda, mengisyaratkan perilaku menyimpang sehubungan dengan jenis pengguna yang mereka nilai. Selanjutnya, menganalisis perilaku Pemfilteran Kolaboratif berbasis pengguna dan Pemfilteran Kolaboratif ketika terus meningkatkan F dengan kenaikan 10%, digambarkan dalam Gambar 2. Kedua bagan mengungkapkan bahwa perbedaan memiliki efek merugikan pada kedua metrik dan pada kedua algoritma Pemfilteran Kolaboratif. Menariknya, kurva presisi dan recall yang sesuai memiliki bentuk yang hampir identik.

Hilangnya akurasi lebih terlihat untuk berbasis item daripada Pemfilteran Kolaboratif berbasis pengguna. Lebih lanjut, baik untuk metrik maupun algoritma Pemfilteran Kolaboratif, penurunan paling mencolok untuk $F \in$

[0,2, 0,4]. Untuk F yang lebih rendah, dampak negatif pada akurasi kecil.

5.3 Eksperimen Online

Eksperimen online membantu berguna untuk memahami implikasi perbedaan topik pada kedua algoritma pemfilteran kolaboratif. Hasil pengamatan diperoleh bahwa efek pendekatan berbeda pada algoritme yang berbeda. Tetapi, dapat mengetahui tentang kekurangan metrik akurasi, penelitian ini menilai kepuasan pengguna yang sebenarnya untuk berbagai tingkat perbedaan, sehingga memerlukan survei online.

Untuk eksperimen online, menghitung setiap jenis daftar rekomendasi baru untuk pengguna dalam kumpulan data yang lebih padat meskipun tanpa K-folding.

5.4 Regresi Linear Berganda

Hasil yang diperoleh dari menganalisis masukan pengguna dengan berbagai sumbu fitur telah menunjukkan bahwa kepuasan pengguna secara keseluruhan dengan daftar rekomendasi tidak hanya bergantung pada keakuratan, tetapi juga pada kisaran minat membaca yang tercakup.

Untuk menilai secara lebih kaku indikasi tersebut melalui metode statistik, menerapkan regresi linier berganda ke hasil survei, memilih nilai daftar keseluruhan sebagai variabel dependen. Sebagai variabel input independen, menyediakan rata-rata suara tunggal dan rentang tertutup, keduanya muncul sebagai polinomial orde pertama dan orde kedua, yaitu masing-masing SVA dan CR, dan SVA2 dan CR2. Juga mencoba beberapa model lain yang lebih kompleks, tanpa mencapai model fitting yang jauh lebih baik.

6. PENELITIAN TERKAIT

Beberapa upaya telah mengatasi masalah dalam membuat daftar teratas-N lebih beragam. Hanya dengan mempertimbangkan literatur tentang pemfilteran kolaboratif dan sistem pemberi rekomendasi secara umum,

tidak ada yang pernah disajikan sebelumnya, sejauh yang diketahui. Penelitian sebelumnya sudah menyelidiki apakah ada efek signifikan dalam keakuratan model prediktif [20].

Namun, beberapa pekerjaan yang terkait dengan pendekatan perbedaan topik dapat ditemukan di pencarian informasi, khususnya mesin pencari meta. Aspek kritis dari desain mesin pencari meta adalah penggabungan beberapa daftar N teratas menjadi satu daftar N teratas. Secara intuitif, daftar N teratas yang digabungkan ini harus mencerminkan peringkat kualitas tertinggi, juga dikenal sebagai "masalah agregasi peringkat" [21]. Kebanyakan pendekatan menggunakan variasi dari model "kombinasi skor linier" (LC), dijelaskan oleh Vogt dan Cottrell [22]. Model LC secara efektif menyerupai skema untuk menggabungkan peringkat asli berbasis akurasi dengan peringkat ketidaksamaan saat ini, tetapi lebih umum dan tidak membahas masalah keragaman. Fagin dkk. [23] mengusulkan metrik untuk mengukur jarak antara daftar N teratas, yaitu, metrik kesamaan antar-daftar, untuk mengevaluasi kualitas peringkat yang digabungkan. Oztekin dkk. Lebih terkait dengan ide dalam membuat daftar yang mewakili keseluruhan minat topik pengguna yang paling banyak, Skema pengelompokan mereka yang mengelompokkan hasil pencarian ke dalam kelompok topik terkait. Pengguna kemudian dapat dengan mudah menelusuri folder topik yang relevan dengan minat pencariannya.

7. KESIMPULAN

Pada perbedaan topik, kerangka algoritma untuk meningkatkan keragaman daftar teratas produk yang direkomendasikan, untuk menunjukkan efisiensinya dalam melakukan perbedaan, juga memperkenalkan metrik intra-daftar kesamaan. Metrik presisi dan perolehan yang kontras, yang dihitung untuk Pemfilteran Kolaboratif berbasis pengguna dan berbasis item serta menampilkan tingkat perbedaan yang berbeda, dengan hasil yang diperoleh dari survei pengguna berskala

besar, hasil menunjukkan bahwa kesukaan pengguna secara keseluruhan terhadap daftar rekomendasi melampaui akurasi dan melibatkan faktor-faktor lain, misalnya, keragaman daftar yang dirasakan pengguna.

Dengan demikian, dapat memberikan bukti empiris bahwa daftar lebih dari sekadar kumpulan rekomendasi tunggal, tetapi mengandung nilai tambah yang intrinsik. Meskipun efek perbedaan sebagian besar marginal pada Pemfilteran Kolaboratif berbasis pengguna, kinerja Pemfilteran Kolaboratif berbasis item meningkat secara signifikan, sebuah indikasi bahwa ada beberapa perbedaan perilaku antara kedua kelas Pemfilteran Kolaboratif. Selain itu, sementara Pemfilteran Kolaboratif berbasis item murni tampak sedikit lebih rendah daripada Pemfilteran Kolaboratif berbasis pengguna murni dalam kepuasan secara keseluruhan, perbedaan Pemfilteran Kolaboratif berbasis item dengan faktor $F_2 [0,3, 0,4]$ membuat Pemfilteran Kolaboratif berbasis item mengungguli Pemfilteran Kolaboratif berbasis pengguna. Menariknya untuk $F_{0.4}$, tidak lebih dari tiga item yang cenderung berubah sehubungan dengan daftar aslinya. Dengan demikian, perubahan kecil memiliki dampak yang tinggi. Hasil penelitian penting untuk skenario aplikasi praktis, karena banyak sistem rekomendasi komersial yang berbasis item, berkat efisiensi komputasi algoritma. Algoritma kerangka kerja dapat meningkatkan keragaman daftar N teratas dari produk yang direkomendasikan, juga telah disampaikan metrik kesamaan intra-daftar. Metrik presisi dan perolehan yang kontras, yang dihitung untuk Pemfilteran Kolaboratif berbasis pengguna dan berbasis item serta menampilkan tingkat perbedaan yang berbeda, dengan hasil yang diperoleh dari survei pengguna berskala besar, menunjukkan bahwa kesukaan pengguna secara keseluruhan terhadap daftar rekomendasi melampaui akurasi dan melibatkan faktor-faktor lain, misalnya, keragaman daftar yang dirasakan pengguna. Dengan demikian,

dapat memberikan bukti empiris bahwa daftar lebih dari sekadar kumpulan rekomendasi tunggal, tetapi mengandung nilai tambah yang intrinsik.

Meskipun efek perbedaan sebagian besar marginal pada Pemfilteran Kolaboratif berbasis pengguna, kinerja Pemfilteran Kolaboratif berbasis item meningkat secara signifikan, sebuah indikasi bahwa ada beberapa perbedaan perilaku antara kedua kelas Pemfilteran Kolaboratif. Selain itu, sementara Pemfilteran Kolaboratif berbasis item murni tampak sedikit lebih rendah daripada Pemfilteran Kolaboratif berbasis pengguna murni dalam kepuasan keseluruhan, perbedaan Pemfilteran Kolaboratif berbasis item dengan faktor $F_2 [0,3, 0,4]$ membuat Pemfilteran Kolaboratif berbasis item mengungguli Pemfilteran Kolaboratif berbasis pengguna. Menariknya untuk $F_{0.4}$, tidak lebih dari tiga item yang cenderung berubah sehubungan dengan daftar aslinya, sehingga perubahan kecil berdampak tinggi.

DAFTAR PUSTAKA

- [1] J. Ben Schafer, J. A. Konstan, and J. Riedl, "Meta-recommendation systems: User-controlled integration of diverse recommendations," in *International Conference on Information and Knowledge Management, Proceedings*, 2002, pp. 43–51.
- [2] J. L. Herlocker, "An algorithmic framework for performing collaborative filtering," *SIGIR '99 Proc. 22nd Annu. Int. ACM SIGIR Conf. Res. Dev. Inf. Retr.*, vol. 7, no. 2, pp. 65–83, 1999.
- [3] L. Greg, "Amazon.com recommendations: item-to-item collaborative filtering," *J. Mark. Res.*, vol. 56, no. 3, pp. 401–418, 2019.
- [4] Z. Munawar, "Keamanan Pada E-Commerce Usaha Kecil dan Menengah," *Temat. - J. Teknol. Inf. dan Komun.*, vol. 5, no. 1 SE-Articles, Jun. 2018.
- [5] C. Hayes, P. Massa, P. Avesani, and P. Cunningham, "An on-line evaluation framework for recommender systems,"

- Recomm. Pers. eCommerce*, no. June, p. 50, 2002.
- [6] N. I. Putri, "Sistem pakar diagnosa tingkat kecanduan gadget pada remaja menggunakan metode Certainty Factor." UIN Sunan Gunung Djati Bandung, 2018.
- [7] B. Sarwar, G. Karypis, J. Konstan, and J. Riedl, "Item-Based Collaborative Filtering Recommendation Algorithms," in *World Wide Web Conference*, 2001, pp. 285–295.
- [8] D. Cosley, S. Lawrence, and D. M. Pennock, "REFeree: An open framework for practical testing of recommender systems using ResearchIndex," in *28th International Conference on Very Large Databases*, 2002, pp. 35–46.
- [9] N. I. Munawar, Zen and Putri, "Keamanan Jaringan Komputer Pada Era Big Data," *J-SIKA| J. Sist. Inf. Karya Anak Bangsa*, vol. 02, no. 01, pp. 14–20, 2020.
- [10] D. M. Nichols, "Implicit Rating and Filtering," in *World Wide Web Internet And Web Information Systems*, 1997, pp. 31–36.
- [11] U. Shardanand and P. Maes, "Social information filtering: Algorithms for automating word of mouth," in *Proceedings of the ACM CHI Conference on Human Factors in Computing Systems*, 1995, pp. 210–217.
- [12] J. Herlocker, Jonathan L.; Konstan, Joseph A.; Borchers, Al; Riedl, "Evaluating Collaborative filtering recommender systems," *ACM Trans. Inf. Syst.*, vol. 22, no. 1, pp. 5–53, 2004.
- [13] B. Sarwar, G. Karypis, J. Konstan, and J. Riedl, "Application of Dimensionality Reduction in Recommender System A Case Study Technical Report CS-TR 00-043," *Comput. Sci. Eng. Dept., Univ. Minnesota*, no. August, 2000.
- [14] and C. K. John S. Breese, David Heckerman, "Empirical Analysis of Predictive Algorithms for Collaborative Filtering," Morgan Kaufmann Publishers Inc, 1998.
- [15] A. I. Schein, A. Popescul, L. H. Ungar, and D. M. Pennock, "Methods and metrics for cold-start recommendations," in *Proceedings of the 25th annual international ACM SIGIR conference on Research and development in information retrieval - SIGIR '02*, 2002, pp. 253–260.
- [16] M. R. McLaughlin and J. L. Herlocker, "A collaborative filtering algorithm and evaluation metric that accurately model the user experience," in *Proceedings of Sheffield SIGIR - Twenty-Seventh Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, 2004, pp. 329–336.
- [17] K. Ali, "TiVo: Making Show Recommendations Using a Distributed Collaborative Filtering Architecture," 2004, pp. 394–401.
- [18] www.bookcrossing.com, "What is BookCrossing?," 2020. [Online]. Available: <http://www.bookcrossing.com>.
- [19] Z. Munawar, B. Siswoyo, and N. S. Herman, "Machine learning approach for analysis of social media," *ADRI Int. Journal. Information. Technol.*, vol. 1, pp. 5–8, 2017.
- [20] Z. Munawar, "Penggunaan Profil Media Sosial Untuk Memprediksi Kepribadian," *Temat. - J. Teknol. Inf. dan Komun.*, vol. 4, no. 2 SE-Articles, Dec. 2017.
- [21] C. Dwork, R. Kumar, M. Naor, and D. Sivakumar, "Rank aggregation methods for the web," in *Proceedings of the 10th International Conference on World Wide Web, WWW 2001*, 2001, pp. 613–622.
- [22] C. C. Vogt and G. W. Cottrell, "Fusion via a linear combination of scores," *Inf. Retr. Boston.*, vol. 1, no. 3, pp. 151–173, 1999.
- [23] R. Fagin, R. Kumar, and D. Sivakumar, "Comparing top k lists," *SIAM J. Discret. Math.*, vol. 17, no. 1, pp. 134–160, 2003.