

BAB II

TINJAUAN PUSTAKA

II.1. Analisis Perbandingan

Analisis perbandingan merupakan suatu proses yang dilakukan untuk membandingkan suatu metode tertentu yang dimana proses ini dilakukan dengan cara mengevaluasi suatu permasalahan atau suatu kondisi dan setelah masalah ditemukan maka dicari jalan keluar untuk memecahkan suatu permasalahan tersebut. Analisis dilakukan untuk menentukan variabel dan alur kerja sistem data mining dalam menentukan pola penjualan pada penjualan sparepart mobil dengan menggunakan data transaksi penjualan kemudian diterapkan dengan menggunakan dua algoritma yaitu algoritma apriori dan algoritma Fp-Growth.(Pratama et al., 2009)

II.2. Data Mining

Data mining merupakan proses penambangan data atau proses pencarian data yang baru dengan mencari pola dan aturan tertentu dari data yang besar. *Data mining* juga di definisikan sebagai knowledge discovery in database (KDD) yaitu kegiatan yang meliputi pengumpulan, pemakaian data, riwayat untuk menemukan keteraturan, pola atau hubungan dalam set data yang berukuran besar. *Data mining* diartikan sebagai proses penemuan pola – pola dalam data.proses ini biasanya otomatis atau bisa juga semiotomatis.pola yang di temukan harus

memiliki arti dan keuntungan, biasanya keuntungan secara ekonomi dan data yang dibutuhkan dalam jumlah yang besar.(Santoso et al., 2016)

Data mining adalah proses ekstraksi informasi dari kumpulan data yang besar melalui penggunaan algoritma dan teknik yang melibatkan bidang ilmu statistik, mesin pembelajaran, dan sistem manajemen *database*. Analisis asosiasi atau *association rule mining* adalah teknik *data mining* untuk menentukan aturan asosiasi antara suatu kombinasi *item* (Ummi, 2016).

II.3. Teknik - Teknik Data Mining

1. Deskripsi

Terkadang peneliti dan analis secara sederhana ingin mencoba mencari cara untuk menggambarkan pola dan kecenderungan yang terdapat dalam data. Deskripsi dari pola dan kecenderungan sering memberikan kemungkinan penjelasan untuk suatu pola atau kecenderungan. Aturan yang dihasilkan harus mudah dimengerti agar dapat dengan efektif meningkatkan tingkat pengetahuan

(*knowledge*) pada sistem. Tugas deskriptif merupakan tugas data mining yang sering dibutuhkan pada teknik *postprocessing* untuk melakukan validasi dan menjelaskan hasil dari proses data mining. *Postprocessing* merupakan proses yang digunakan untuk memastikan hanya hasil yang valid dan berguna yang dapat digunakan oleh pihak yang berkepentingan.

2. Estimasi

Estimasi hampir sama dengan klasifikasi, kecuali variabel target estimasi lebih ke arah numerik daripada ke arah kategori. Model dibangun menggunakan record lengkap yang menyediakan nilai dari variabel target sebagai nilai prediksi.

Selanjutnya, pada peninjauan berikutnya estimasi nilai dari variabel target dibuat berdasarkan nilai variabel prediksi. Sebagai contoh akan dilakukan estimasi tekanan darah sistolik pada pasien rumah sakit berdasarkan umur pasien, jenis kelamin, indeks berat badan, dan level sodium darah. Hubungan antara tekanan darah sistolik dan nilai variabel prediksi dalam proses pembelajaran akan menghasilkan model estimasi. Model estimasi yang dihasilkan dapat digunakan untuk kasus baru lainnya

3. Prediksi

Prediksi hampir sama dengan klasifikasi dan estimasi, kecuali bahwa dalam prediksi nilai dari hasil akan ada dimasa mendatang. Beberapa metode dan teknik yang digunakan dalam klasifikasi dan estimasi dapat pula digunakan (untuk keadaan yang tepat) untuk prediksi. Contoh dari tugas prediksi misalnya

- a. untuk memprediksikan adanya pengurangan jumlah pelanggan dalam waktu dekat.
- b. prediksi harga saham dalam tiga bulan yang akan datang

4. Klasifikasi

Dalam klasifikasi, terdapat target variabel kategori. Sebagai contoh, penggolongan pendapatan dapat dipisahkan dalam tiga kategori, yaitu pendapatan tinggi, pendapatan sedang, dan pendapatan rendah. Contoh lain klasifikasi dalam bisnis dan penelitian adalah:

- a. Menentukan apakah suatu transaksi kartu kredit merupakan transaksi yang curang atau tidak.

- b. Memperkirakan apakah suatu pengajuan hipotek oleh nasabah merupakan suatu kredit yang baik atau buruk.
- c. Mendiagnosis penyakit seorang pasien untuk mendapatkan termasuk kategori penyakit apa.

5. Pengklusteran

Pengklusteran merupakan pengelompokan record, pengamatan, atau memperhatikan dan membentuk kelas objek - objek yang memiliki kemiripan. Kluster adalah kumpulan record yang memiliki kemiripan satu dengan yang lainnya dan memiliki ketidakmiripan dengan record-record dalam kluster lain. Tujuannya adalah untuk menghasilkan pengelompokan objek yang mirip satu sama lain dalam kelompok - kelompok. Semakin besar kemiripan objek dalam suatu cluster dan semakin besar perbedaan tiap *cluster* maka kualitas analisis cluster semakin baik.

6. Asosiasi

Tugas asosiasi dalam *Data Mining* adalah menemukan atribut yang muncul dalam satu waktu. Dalam dunia bisnis lebih umum disebut analisis keranjang belanja. *Data Mining* dimaksudkan untuk menangani data dalam ukuran yang demikian, maka salah satu arah penelitian di bidang *Data Mining* adalah mengembangkan algoritma-algoritma tersebut agar dapat menangani data yang berukuran sangat besar. Aturan asosiasi juga sering dinamakan *market basket analysis* (analisis keranjang belanja), aturan asosiasi ingin memberikan informasi dalam bentuk hubungan “*if-then*” atau “jika-maka”. Aturan ini dihitung dari data yang sifatnya probabilitas.

II.4. Algoritma Apriori

Algoritma apriori termasuk jenis aturan asosiasi pada Data Mining. Aturan yang menyatakan asosiasi antara beberapa atribut sering disebut affinity analysis atau market basket analysis. Analisis asosiasi atau association rule mining adalah teknik Data Mining untuk menemukan aturan suatu kombinasi item. Salah satu tahap analisis asosiasi yang menarik perhatian banyak peneliti untuk menghasilkan algoritma yang efisien adalah analisis pola frekuensi tinggi (frequent pattern mining) (Ummi, 2016).

Algoritma apriori adalah algoritma paling terkenal untuk menemukan pola frekuensi tinggi. Pola frekuensi tinggi adalah pola-pola item di dalam suatu database yang memiliki frekuensi atau support di atas ambang batas tertentu yang disebut dengan istilah minimum support. Algoritma apriori dibagi menjadi beberapa tahap yang disebut iterasi atau pass yaitu:

- a. Pembentukan kandidat itemset , kandidat k-itemset dibentuk dari kombinasi (k-1)-itemset yang didapat dari iterasi sebelumnya. Satu cara dari algoritma apriori adalah adanya pemangkasan kandidat k-itemset

- b. yang subset-nya yang berisi k-1 item tidak termasuk dalam pola frekuensi tinggi dengan panjang k-1.

Iterasi 1 mulai dilakukan dengan membentuk kandidat 1 – *itemset* (C1) dari data-data transaksi tersebut dan hitung jumlah *support*-nya. Cara menghitung *support* adalah jumlah kemunculan item dalam transaksi dibagi dengan jumlah seluruh transaksi.

$$Support (A) = \frac{\text{Jumlah Transaksi Mengandung A}}{\text{Total Transaksi}} \times 100\%$$

- c. Penghitungan *support* dari tiap kandidat k-itemset. *Support* dari tiap kandidat k-itemset didapat dengan menscan database untuk menghitung jumlah transaksi yang memuat semua item di dalam kandidat k-itemset tersebut. Ini adalah juga ciri dari algoritma apriori dimana diperlukan penghitungan dengan scan seluruh database sebanyak k - itemset terpanjang. Proses pembentukan C2 atau disebut juga dengan 2 itemset dengan jumlah *minimum support* 10% dapat diselesaikan dengan rumus sebagai berikut :

$$Support (A) = \frac{\text{Jumlah Transaksi Mengandung A}}{\text{Total Transaksi}} \times 100\%$$

Setelah semua pola frekuensi tinggi ditemukan , barulah dicari aturan asosiasi yang memenuhi syarat minimum untuk *confidence* aturan asosiatif **A→B** Minimum *Confidence* = 50% Nilai *Confidence* dari aturan **A→B** diperoleh

$$Confidence P(B|A) = \frac{\sum \text{Total transaksi mengandung A dan B}}{\sum \text{Transaksi mengandung A}} \times 100\%$$

II.5. FP-Growth

Algoritma *FP-Growth* merupakan pengembangan dari algoritma *Apriori*. Sehingga kekurangan dari algoritma *Apriori* diperbaiki oleh algoritma *FP-Growth*. *Frequent Pattern Growth (FP-Growth)* adalah salah satu alternatif algoritma yang dapat digunakan untuk menentukan himpunan data yang paling sering muncul (*frequent itemset*) dalam sebuah kumpulan data. Pada algoritma *Apriori* diperlukan *generate candidate* untuk mendapatkan *frequent itemsets*. Akan tetapi, di algoritma *FP-Growth* *generate candidate* tidak dilakukan karena *FP-Growth* menggunakan konsep pembangunan *tree* dalam pencarian *frequent itemsets*. Hal tersebutlah yang menyebabkan algoritma *FP-Growth* lebih cepat dari algoritma *Apriori*. Karakteristik algoritma *FP-Growth* adalah struktur data yang digunakan berbentuk *tree* yang disebut dengan *FP-Tree*. Dengan menggunakan *FP-Tree*, algoritma *FP-Growth* dapat langsung mengekstrak *frequent itemset* dari *FP-Tree*. (Lintang, dkk, 2017 : 128).

Frequent Pattern Growth (FP-Growth) adalah salah satu alternatif algoritma yang dapat digunakan untuk menentukan himpunan data yang paling sering muncul (*frequent itemset*) dalam sebuah kumpulan data. Algoritma *FP-Growth* merupakan pengembangan dari algoritma *Apriori*. Sehingga kekurangan dari algoritma *Apriori* diperbaiki oleh algoritma *FP-Growth*. *FP-Growth* menggunakan konsep pembangunan *tree* dalam pencarian *frequent itemsets*. Hal tersebutlah yang menyebabkan algoritma *FP-Growth* lebih cepat dari algoritma *Apriori*. Karakteristik algoritma *FP-Growth* adalah struktur data yang digunakan adalah *tree* yang disebut dengan *FP-Tree*. Dengan menggunakan *FP-Tree*, algoritma *FP-Growth* dapat langsung mengekstrak *frequent itemset* dari *FP-Tree*. (Fajrin dan Mulana, 2018 : 30).

$$\text{Support A} = \frac{\text{Jumlah Transaksi Mengandung A}}{\text{Total Transaksi}} \dots\dots\dots(1)$$

$$\text{Confidence} = \frac{\text{Jumlah Transaksi}}{\text{Jumlah Transaksi A Dan B}} \dots\dots\dots(2)$$

Untuk mencari *frequent itemset* dan *association rule* menggunakan algoritma *FP-Growth* ada beberapa langkah-langkah sebagai berikut :

1. Hitung terlebih dahulu *minimum support*.
2. Pembuatan *header item*.
3. Prioritas *item* yang memiliki frekuensi terbesar.
4. Susun *item* berdasarkan prioritas.

5. Bentuk *FP-Tree* dari tiap transaksi
6. Validasi. (Abdullah, 2018 : 24).

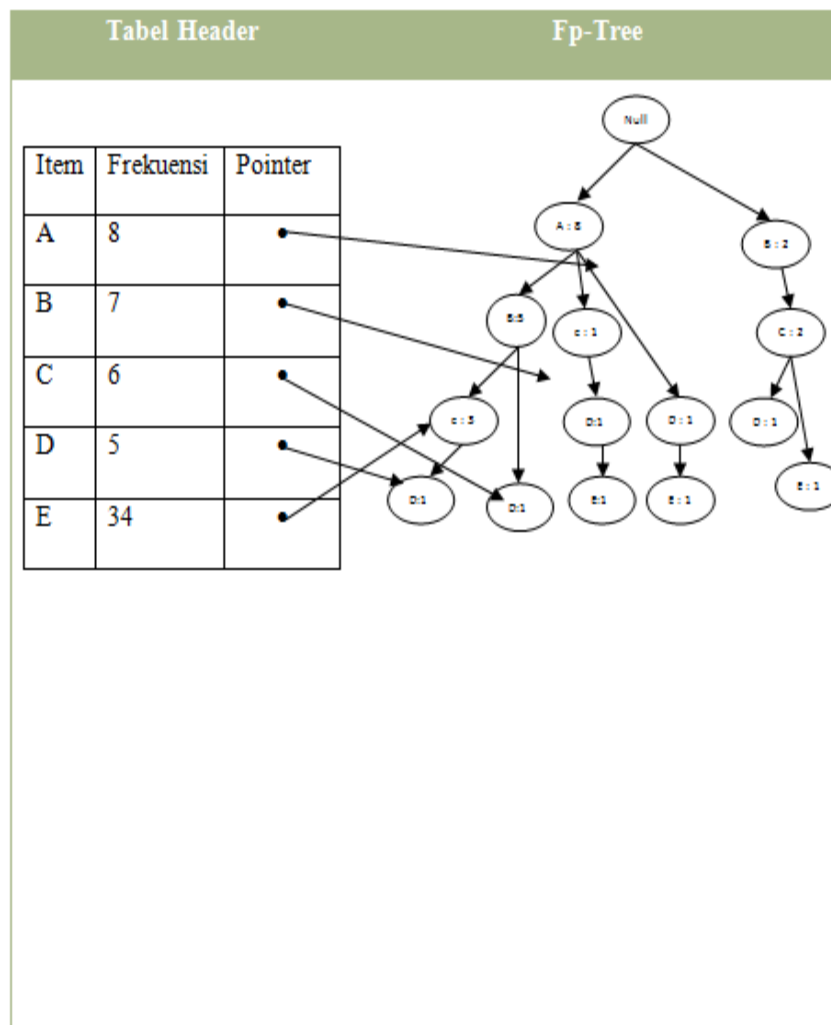
II.6. Fp-Tree

Fp-Tree merupakan suatu struktur penyimpanan yang berbentuk pohon yang dimana penyimpanan dimampatkan. *FpTree* dibangun dengan memetakan setiap data transaksi ke dalam setiap lintasan tertentu dalam *Fp-Tree*. Karena dalam setiap transaksi yang dipetakan, kemungkinan memiliki *item* yang sama, maka lintasannya memungkinkan akan saling menimpa antara satu sama lain. (Gunadi & Sensuse, 2012) Karena semakin banyak data transaksi yang memiliki *item* yang sama, maka proses pemampatan dengan struktur data *Fp-Tree* semakin efektif :

Adapun *Fp-Tree* merupakan sebuah pohon yang memiliki definisi yaitu :

- a. *Fp-Tree* dibentuk oleh sebuah akar yang diberi label *null*, sekumpulan *sub-tree* yang beranggotakan *item-item* tertentu, dan sebuah tabel *frequent header*.

- b. Setiap simpul yang terdapat dalam Fp-Tree Mengandung tiga informasi penting, yaitu menginformasikan jenis *item* yang direpresentasikan simpul tersebut, *support count*, merepresentasikan jumlah lintasan transaksi yang melalui simple tersebut, dan *pointer* penghubung yang menghubungkan simpul – simpul dengan label *item* yang sama antar lintasan yang ditandai dengan garis panah putus – putus.



Gambar 2.1 Aliran Informasi dalam Fp-Tree

(Sumber : (Gunadi & Sensuse, 2012))

II.7. KDD (Knowledge Discovery In Database)

Menurut (Hendrickx et al., 2015) Istilah *data mining* dan *knowledge discovery* adalah proses penggalian informasi tersembunyi dalam suatu basis data yang besar. Sebenarnya kedua istilah tersebut memiliki konsep yang berbeda, tetapi berkaitan satu sama lain. Dan salah satu tahapan dalam keseluruhan proses KDD adalah *data mining*. Proses KDD secara garis besar dapat dijelaskan sebagai berikut :

Adapun tahap - tahap KDD dari proses *Data Mining* adalah sebagai berikut:

1. Seleksi data (*data selection*)

Pada proses ini dilakukan proses pemilihan (Seleksi) data dari sekumpulan data yang besar yang dimana diperlukannya proses penggalian informasi baru .hasil data yang telah diseleksi dipisahkan dengan data operasionalnya kedalam suatu berkas penyimpanan yang baru untuk dilakukannya proses *data mining*.

2. Pembersihan data (*data cleaning*)

Sebelum di lakukannya proses *data mining* maka perlu dilakukannya proses pembersihan data yang dimana Merupakan proses menghilangkan *noise* dan data yang tidak konsisten, serta membuang *duplikasi data* dan memperbaiki kesalahan pada data.

3. *Data transformation*

Data diubah atau digabung ke dalam format yang sesuai untuk diproses dalam *data mining*.Biasanya data dapat berupa *coding*

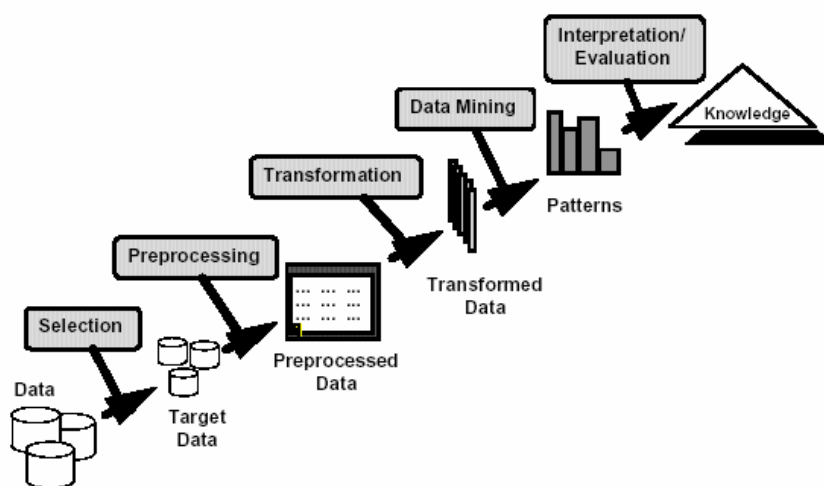
yang dimana data tersebut dapat mempermudah ketika data dibaca oleh *database*.

4. Proses *mining*.

Merupakan suatu proses utama saat metode diterapkan untuk menemukan pengetahuan berharga dan tersembunyi dari data.

5. *Evaluation*

Pola informasi yang dihasilkan dari proses *data mining* perlu ditampilkan dalam bentuk yang mudah dimengerti oleh pihak yang berkepentingan. Tahap ini merupakan bagian dari proses KDD yang disebut *interpretation*. Tahap ini mencakup pemeriksaan apakah pola atau informasi yang ditemukan bertentangan dengan fakta atau hipotesis yang ada sebelumnya.



Gambar 2.2 Aliran Informasi dalam *data mining*

(Sumber : (Lumbantoruan, 2015))

II.8. Sparepart

Sparepart sendiri merupakan suatu komponen yang membentuk suatu kesatuan yang setiap komponennya memiliki fungsi tertentu. setiap mobil memiliki kebutuhan akan suku cadang atau *sparepart* yang berbeda – beda dan memiliki tipe – tipe yang berbeda pula. Ada beberapa komponen yang harus diperhatikan yaitu seperti Radiator, Ban, Rem, Transmisi, Suspensi Alternator, Powertrain Control Module (PCM), Kolom Kemudi, Dan Fuel Injection. Beberapa *sparepart* ini sangat berpengaruh pada kendaraan yang di gunakan oleh pengemudi karena apabila salah satu *sparepart* terdapat kerusakan dapat berdampak sangat berbahaya bagi pengemudi (<http://hakimsimanjuntak.blogspot.com/2010/11/pengertian-sukucadang-spare-part.html/11Mei2020>)

Secara umum *sparepart* dapat di bagi menjadi 2 jenis yaitu :

1. *Sparepart* Baru yaitu salah satu komponen yang memiliki kondisi baru yang belum pernah di pakai sama sekali.
2. *Sparepart* Bekas yaitu salah satu komponen yang telah di pakai dengan kondisi masih layak pakai sehingga aman untuk di gunakan oleh pengemudi.

II.9. RapidMiner

Rapid Miner adalah suatu alat atau perangkat lunak yang bersifat terbuka (open source). Rapid Miner adalah solusi untuk penambangan data analitik, penambangan teks, dan analisis prediktif. Rapid Miner menggunakan berbagai teknik deskriptif dan prediktif untuk memberikan wawasan kepada

pengguna sehingga mereka dapat membuat keputusan terbaik. Rapid Miner memiliki sekitar 500 operator data mining, termasuk operator input, output, preprocessing data, dan visualisasi. Rapid Miner adalah perangkat lunak independen untuk analisis data dan mesin penambangan data yang dapat diintegrasikan ke dalam produknya sendiri. Rapid

Miner ditulis dalam bahasa Java, sehingga dapat bekerja di semua sistem operasi. Nama RapidMiner sebelumnya adalah YALE (sampai saat ini Lingkungan belajar lain), versi aslinya dikembangkan pada tahun 2001 oleh Ralf Klinkenberg, Ingo Mierswa dan Simon Fischer dari Departemen Kecerdasan Buatan Universitas Dortmund. Rapid Miner dirilis di bawah lisensi AGPL (GNU Affero General Public License) versi 3. Sejauh ini, ribuan aplikasi telah dikembangkan menggunakan Rapid Miner di lebih dari 40 negara. Tidak ada keraguan bahwa Rapid Miner adalah software open source untuk data mining, karena software tersebut sudah menjadi pemimpin dunia. Dalam jajak pendapat yang dilakukan oleh portal data mining KDnuggets dari 2010 hingga 2011, Rapid Miner dinilai sebagai perangkat lunak penambangan data nomor satu. RapidMiner menyediakan GUI (pengguna grafis Interface) untuk merancang dan menganalisis jalur pipa. GUI akan menghasilkan file XML (Extensible Markup Language) yang mendefinisikan proses analisis yang ingin diteruskan pengguna ke data. RapidMiner akan membaca file ini untuk menjalankan penganalisis secara otomatis (Aprilia et al., 2013).

Tabel II.1 Penelitian Terkait dengan Algoritma *Apriori* dan Algoritma *Fp-Growth*

No	Nama Penulis	Publish	Judul	Pembahasan
1	Domi Sepri, M.Afdal	Jurnal Sistem Informasi Kaputama (JSIK), Vol 1 No 1, Januari 2017	Analisa Dan Perbandingan Metode Algoritma Apriori Dan Fp-Growth Untuk Mencari Pola Daerah Strategis Pengenalan Kampus Studi	Dalam penelitian ini penulis mencoba membandingkan hasil dari algoritma Apriori dan FP-Growth yang menggunakan data mahasiswa angkatan 2015/2016 dengan nilai minsupport = 0.05% dan nilai minconfidence = 0.7% telah diperoleh 19 Association Rule dan 2 rule tertinggi yang dapat dijadikan sebagai pengetahuan baru serta acuan berharga pada lingkup penelitian ini. kesimpulan dari analisa penerapan pada algoritma FP-Growth maupun Apriori terhadap knowledge yang baru didapat yaitu berupa penentuan dalam menentukan strategi dalam mempromosikan pendidikan.

No	Nama Penulis	Publish	Judul	Pembahasan
			Kasus Di Stkip Adzkia Padang	
2	Rizky Fitria,Warnia Nengsih dan Dini Hidayatul Qudsi	Jurnal Sistem Informasi (Journal of Information Systems). 2/13 (2017), 118- 124	Implementasi Algoritma Fp- Growth Dalam Penentuan Pola Hubungan Kecelakaan Lalu Lintas	Berdasarkan uraian yang telah dijelaskan, maka dapat diambil kesimpulan bahwa aplikasi penentuan pola kecelakaan lalu lintas dengan algoritma Fpgrowth yang dibangun dapat berjalan dengan baik dan sesuai dengan tujuan yang diharapkan. Meskipun demikian, aplikasi ini hanya membantu pihak terkait Dalam membuat keputusan untuk meminimalisir terjadinya kecelakaan, bukan untuk membuat keputusan dalam membuat kebijakan. Berdasarkan pengujian nilai minimum support dan pengujian nilai minimum confidence didapatkan

No	Nama Penulis	Publish	Judul	Pembahasan
				nilai minimum support yang tepat yaitu 0.04 dan nilai confidence yang tepat adalah 0.6. Besar kecilnya nilai minimum support dan confidence ini nantinya akan mempengaruhi jumlah pola yang dihasilkan

No	Nama Penulis	Publish	Judul	Pembahasan
3	Agus Junaidi	Jurnal SISFOKOM, Volume 08, Nomor 01, Maret 2019	Implementasi Algoritma Apriori dan FP-Growth Untuk Menentukan Persediaan Barang	Peletakan barang yang sesuai dengan hubungan antar barang yang biasanya dibeli konsumen juga dapat ditentukan berdasarkan support dan confidence berdasarkan market base analysis yang diperoleh dari perhitungan asosiasi dengan menggunakan metode Frequent Pattern Growth. Dengan menggunakan metode Frequent Pattern Growth maka penempatan barang dan persediaan barang pada minimarket dapat terkontrol dengan baik sehingga pelayanan pada konsumen akan meningkat yang akhirnya dapat juga meningkatkan omset penjualan. Dalam penelitian ini support ditentukan menggunakan ambang batas 60% dan confidence 90%. Secara keseluruhan dari data sampel penjualan diperoleh 152 rule yang terdiri dari 24 rule asosiasi yang memenuhi support dengan ambang batas 60% dan 108 rule yang memenuhi confidence 90%.

No	Nama Penulis	Publish	Judul	Pembahasan
4	Alhassan Bala, Mansur Zakariyya Shuaibu,Zahara ddeen Karami Lawal, Rufa'i Yusuf Zakari	Alhassan Bala et al, Int.J.Computer Technology & Applications,V ol 7(2),279- 293	Performance Analysis of Apriori and FP-Growth Algorithms (Association Rule Mining)	Makalah ini berfokus pada perbandingan dan analisis penambangan aturan asosiasi berdasarkan parameter tertentu yang meliputi jumlah contoh, tingkat kepercayaan yang berbeda dan Tentu saja tingkat dukungan yang berbeda diterapkan di Weka mengambil set data supermarket sebagai contoh data. Dua aturan asosiasi teknik yang digunakan termasuk FP-growth Algoritma dan algoritma Apriori. Saya tersebut jelas bahwa, FP-growth Algoritma lebih cepat dalam hal eksekusi waktu dibandingkan dengan Algoritma Apriori. Ini menunjukkan bahwa jumlah waktu yang dibutuhkan untuk menjalankan untuk menyelesaikannya kurang dari jumlah waktu dibutuhkan oleh algoritma Apriori.

No	Nama Penulis	Publish	Judul	Pembahasan
5	Mustakim Folkes E. Laumal, Della Maulina Herianda, Ahmad Ilham, Achmad Daeng GS, Ida Bagus Ary Indra Iswara, Nuning Kurniasih, Akbar	IOP Conf. Series: Journal of Physics: Conf. Series 1114 (2018) 012131	Market Basket Analysis Using Apriori and FP- Growth for Analysis Consumer Expenditure Patterns at Berkah Mart in Pekanbaru Riau	Perbandingan waktu minimum digunakan untuk membentuk aturan dari setiap algoritma yang memiliki perbedaan waktu yang sangat lama. Itu direkam dalam hitungan detik; Algoritma FP-Growth hanya membutuhkan 25% dari waktu yang dibutuhkan untuk memproses 1 iterasi dalam sebuah aturan, sedangkan Apriori menggunakan waktu minimum 75% untuk 1 proses pembentukan aturan. Hasil percobaan menunjukkan bahwa penerapan algoritma FP-Growth dan Apriori untuk analisis belanja konsumen Pola di Berkah Mart dapat meningkatkan pendapatan atau keuntungan secara keseluruhan, namun disarankan menggunakan FPGrowth algoritma yang secara maksimal memiliki kecepatan proses dalam bentuk aturan dan memiliki dukungan dan nilai kepercayaan algoritma apriori yang unggul.

No	Nama Penulis	Publish	Judul	Pembahasan
	Iskandar, Gloria Manulanggaan d Robbi Rahim			
6	Asrul Abdullah	Jurnal Ilmu Komputer dan Informatika ISSN: 2621- 038X, Online ISSN: 2477- 698X	Rekomendasi Paket Produk Guna Meningkatkan Penjualan Dengan Metode FP-Growth	Aplikasi rekomendasi paket produk menggunakan metode FP-Growth dapat menjadi alternatif bagi para penentu keputusan/penjual untuk memilih produk yang bisa digabungkan di dalam satu paket yang diharapkan berdampak pada peningkatan penjualan. FP tree memerlukan dua kali scanning database untuk menentukan frequent itemset sehingga membuatnya lebih efisien dibandingkan Apriori. Hasil dari association rules menggunakan FP-Growth dijadikan rekomendasi bagi para

No	Nama Penulis	Publish	Judul	Pembahasan
				penjual/retailer dalam memberikan paket penjualan barang bagi konsumennya dan hasil dari penelitian ini adalah ditemukan dua pasangan item barang yakni kopi, gula dan teh, susu yang memiliki support sebesar 30% dan confidence sebesar 70%.
7	Desti Fitriati,Musi Hardiyanto	Prosiding SNATIF Ke -5 Tahun 2018	Perbandingan Algoritma Apriori Dan Algoritma Fp- Growth Untuk Mengetahui Pola Penggunaan Transportasi	Berdasarkan hasil pengolahan data kuesioner menggunakan algoritma apriori dan algoritma fpgrowth didapatkan hasil akurasi yang sama dengan kombinasi itemset pada algoritma apriori jika masyarakat memilih transportasi online dengan alasan “transportasi motor (D1)” maka masyarakat memilih dengan alasan “tidak pernah memiliki history driver” dengan akurasi sebesar 86%. Sedangkan hasil uji menggunakan algoritma fpgrowth didapatkan kombinasi itemset jika masyarakat memilih transportasi online dengan alasan “menaiki

No	Nama Penulis	Publish	Judul	Pembahasan
			Online	transportasi online dengan sendirian (G1)” maka masyarakat memilih dengan alasan “transportasi motor (D1)” dengan akurasi sebesar 86%. Jadi pengolahan data mining menggunakan algoritma apriori dan algoritma fpgrowth memiliki tingkat akurasi yang baik untuk meningkatkan pelayanan perusahaan transportasi online kepada masyarakat dan tentunya dapat meningkatkan laba perusahaan dengan pelayanan yang terus ditingkatkan.
8	Dio Prima Mulya	Jurnal Teknologi Dan Sistem Informasi Bisnis Vol. 1 No. 1 Januari	Analisa Dan Implementasi Association Rule Dengan Algoritma Fp-Growth Dalam Seleksi	Dengan Algoritma FP-Growth dapat membantu seleksi pembelian tanah liat dengan menghasilkan rule. Algoritma FP-Growth memberikan informasi eksklusif dan menggambarkan proses dalam pembelian tanah liat oleh perusahaan. Penggunaan Algoritma Fp-Growth dan aplikasi software Rapidminer 7.4.0 didapatkan hasil berupa aturan (rules) yang merupakan kumpulan dari beberapa frequent itemset dengan nilai confidence yang tinggi.

No	Nama Penulis	Publish	Judul	Pembahasan
		2019	Pembelian Tanah Liat (Studi Kasus Di Pt. Anveve Ismi Berjaya)	Hasil rule dari Algoritma FP-Growth dan software Rapidminer Studio 7.4.0 membantu untuk pimpinan perusahaan membuat suatu perencanaan pembelian tanah liat.
9	Nugroho Wandu, Rully A. Hendrawan, dan Ahmad Mukhlason	Jurnal Teknik Its Vol. 1, (Sept, 2012) Issn: 2301- 9271	Pengembangan Sistem Rekomendasi Penelusuran Buku dengan Penggalan Association Rule Menggunakan	penggunaan nilai minsup yang semakin besar akan memberikan rekomendasi berdasarkan buku-buku yang semakin sering dipinjam dalam kumpulan transaksi, dan sebaliknya, semakin kecil nilai minsup yang digunakan, artinya rekomendasi yang diberikan berasal dari buku yang semakin jarang dipinjam. Untuk minconf, semakin besar nilainya, semakin besar kemungkinan muncul buku yang direkomendasikan ketika anggota memilih buku tertentu. Sebaliknya, semakin kecil nilai minconf, maka beberapa rekomendasi yang

No	Nama Penulis	Publish	Judul	Pembahasan
			Algoritma Apriori (Studi Kasus Badan Perpustakaan dan Kearsipan Provinsi Jawa Timur)	tidak signifikan akan ditemukan dan digolongkan ke dalam rekomendasi. Pengukuran validitas algoritma menggunakan dua scenario dan menunjukkan hasil yang sama dalam perhitungannya, sehingga aplikasi dapat dikatakan valid.
10	Daniel Hunyadi	Proceedings of the European Computing Conference	Performance comparison of Apriori and FP- Growth algorithms in generating	lgoritme paling umum yang digunakan untuk jenis tindakan ini adalah Apriori (yang menghasilkan kumpulan sering dan aturan asosiasi) dan Aturan Asosiasi FP-Pertumbuhan / Buat (FP-Pertumbuhan menghasilkan kumpulan artikel yang sering, yang kemudian digunakan oleh Buat Aturan Asosiasi yang akan dibuat

No	Nama Penulis	Publish	Judul	Pembahasan
			association rules	aturan asosiasi). Meskipun algoritma Apriori memproses data dengan cara yang berbeda dari algoritma FPGrowth dan Create Association Rules, menghilangkan kumpulan artikel yang tidak sering (dengan dukungan minimum lebih kecil dari dukungan minimum yang ditentukan), terdapat korelasi yang signifikan ($p < 0,05$) antara hasil proses yang dihasilkan melalui algoritme masing-masing, yang dibuktikan melalui garis regresi, dalam hal dukungan independen, masing-masing melalui garis regresi, dalam kasus ini dari tiga varian nilai min_support.
11	Ravi Ranjan, Aditi Sharma	4th International Conference on Computers and	Evaluation of Frequent Itemset Mining Platforms using Apriori and	Studi dalam karya ini menyajikan MR Apriori berbasis Flink dan Pertumbuhan FP Paralel yang diterapkan pada pola-pola yang sering menambang dari ekstensif kumpulan data. Ini menggunakan persyaratan Apriori penting bahwa itemset harus sering jika hanya semua subset yang

No	Nama Penulis	Publish	Judul	Pembahasan
		Management	FPGrowth Algorithm	tidak kosong sering. Ini dijalankan pada Apache Flink, Apache Spark, dan Apache Hadoop yang memberikan kondisi pemrosesan paralel dan terdistribusi. Flink paling sesuai untuk Apriori karena Apache Flink memiliki dukungan lokal untuk kalkulasi iteratif dan Apriori didasarkan pada kalkulasi iteratif. Selain itu, peneliti dapat menerapkan versi paralel dari Apriori dan FP Growth ke berbagai domain aplikasi seperti data cuaca, lalu lintas internet, informasi medis, dll. Kami juga dapat menggunakan algoritme ini untuk menghasilkan aturan asosiasi yang berbeda dan menarik dengan lebih cepat dan efektif.
12	Islamiyah, Putri Lestari Ginting, Nataniel	The 6th International Conference on	Comparison of Priori and FP- Growth Algorithms	Meskipun jumlah aturan asosiasi yang dihasilkan berbeda, aturan asosiasi ini memiliki nilai asosiasi akhir yang sama ($\text{dukungan} \times \text{keyakinan}$). Tiga nilai asosiasi tertinggi terakhir dalam urutan tersebut adalah 2,38% jika Anda

No	Nama Penulis	Publish	Judul	Pembahasan
	Dengen, Medi Taruk	Electrical, Electronics and Information Engineering (ICEEIE 2019)	in Determining Association Rules	membeli Biskuit, Anda akan membeli Snack; 2,35% jika Anda membeli Snack, Anda akan membeli Biskuit; dan 2,09% jika Anda membeli Snack, Anda akan membeli Susu. Waktu eksekusi yang dibutuhkan untuk mengeksekusi 979 kategori item pada aplikasi menyatakan bahwa algoritma FP-Growth lebih cepat dari pada algoritma Apriori. Semakin besar dukungan minimum yang ditentukan, semakin sedikit Aturan Asosiasi yang dibentuk.
13	K.Dharmaraajan, M.A. Dorairangaswamy	2016 IEEE International Conference on Advances in Computer Applications	Analysis of FP-Growth and Apriori Algorithms on Pattern Discovery from Weblog Data	Untuk memindai pola Sering yang besar, file Algoritma Apriori membutuhkan interval yang lebih lama jika dibandingkan dengan algoritma lain. Pohon pola frekuensi (FP-tree) digunakan untuk menyimpan informasi penting yang padat tentang frekuensi pola. FP-Growth untuk penambangan yang efektif dari pola yang sering dalam kumpulan data yang sangat besar. Algoritma FP-Growth yang digunakan

No	Nama Penulis	Publish	Judul	Pembahasan
		(ICACA)		dalam data weblog adalah lebih efektif daripada algoritma Apriori. Juga membutuhkan lebih sedikit waktu dan kinerjanya lebih baik. Pekerjaan penelitian ini membantu meningkatkan industri kelistrikan untuk mempromosikan produk. Ini metodologi dapat diperluas ke industri berbasis online lainnya mempertahankan pelanggan dan meningkatkan pendapatan perusahaan.
14	Ariefana Ria Riszky, Mujiono Sadikin	Jurnal Teknologi dan Sistem Komputer, 7(3), 2019, 103-108	Data Mining Menggunakan Algoritma Apriori untuk Rekomendasi Produk	Algoritma apriori yang diujicobakan pada dataset transaksi penjualan fleksibel untuk digunakan sebagai dasar pengambilan keputusan perusahaan pada area pemasaran. Aturan asosiasi yang terbentuk dapat digunakan sebagai acuan untuk rekomendasi produk yang memenuhi nilai confidence dan support minimum.

No	Nama Penulis	Publish	Judul	Pembahasan
			bagi Pelanggan	
15	Nur Fitrianti Fahrudin	MIND Journal ISSN (p): 2528-0015 ISSN (e): 2528-0902	Penerapan Algoritma Apriori untuk Market Basket Analysis	Untuk menentukan aturan asosiasi yang diambil, setelah melakukan analisis asosiasi, pada tahap akhir diperlukan perhitungan lift ratio untuk melihat kekuatan dari aturan asosiasi yang telah dibentuk. Berdasarkan hasil pengujian dengan menghitung lift ratio didapatkan beberapa aturan yang memiliki nilai ratio tinggi yaitu : {Mentega} $\square$ {telur} dengan nilai confidence 1 dan lift ratio 1.25, {telur} -> {terigu} dengan nilai confidence 75% dan lift ratio 1.25, {gula,terigu} -> {telur} dengan nilai confidence 100% dan lift ratio 1.25, terakhir {telur,terigu} -> {gula} dengan nilai confidence 100% dan lift ratio 1.25.Seperti yang telah dicontohkan pada item {mentega} -> {telur}, manajer toko dapat mengambil strategi untuk menyimpan mentega di dekat telur untuk meningkatkan penjualan. Selain menyimpan metega di dekat telur, manajer

No	Nama Penulis	Publish	Judul	Pembahasan
				toko juga dapat memberikan diskon untuk pembelian kedua item tersebut.