

PERANCANGAN MODEL SISTEM INTELIJENSIA BISNIS



Business Intelligence (BI) atau Intelijensia Bisnis didefinisikan sebagai model matematika dan analisa metodologi eksploitasi data yang tersedia untuk mendapatkan informasi dan pengetahuan yang berguna dalam proses pengambilan keputusan yang kompleks. Intelijensia Bisnis adalah kombinasi dari data *warehouse* dan sistem intelijen. Sistem intelijensia bisnis dapat berperan serta sebagai alat untuk memberikan informasi yang akurat dan berguna bagi pengambil keputusan dalam batas waktu yang ditentukan untuk mendukung pengambilan keputusan.

Intelijensia bisnis menggabungkan analisis bisnis, penambangan data, visualisasi data, alat dan infrastruktur data, serta praktik terbaik untuk membantu organisasi membuat keputusan yang lebih berdasarkan data. Intelijensia bisnis modern memiliki pandangan komprehensif tentang data organisasi dan menggunakan data tersebut untuk mendorong perubahan, menghilangkan inefisiensi, dan dengan cepat beradaptasi terhadap perubahan pasar atau pasokan. Solusi BI modern memprioritaskan analisis layanan mandiri yang fleksibel, data yang diatur pada platform tepercaya, pemberdayaan pengguna bisnis, dan kecepatan dalam memperoleh wawasan.



PENERBIT WAWASAN ILMU
Anggota IKAPI (215/JTE/2021)

Email : redaksi@wawasanilmu.co.id
WA : 089 535 969 2310
FB : Penerbit Wawasan Ilmu
IG : @penerbitwawasanilmu
@katalogwawasanilmu
Web : www.wawasanilmu.co.id



Rina Fitriana | Dedy Sugiarto
Miwan Kurniawan Hidayat | Yusri Eli Hotman Turnip

PERANCANGAN MODEL SISTEM INTELIJENSIA BISNIS

Rina Fitriana
Dedy Sugiarto
Miwan Kurniawan Hidayat
Yusri Eli Hotman Turnip



PERANCANGAN MODEL SISTEM INTELIJENSIA BISNIS



Perancangan Model Sistem Intelijensia Bisnis

Undang-Undang Republik Indonesia Nomor 28 Tahun 2014 tentang Hak Cipta

Fungsi dan sifat hak cipta Pasal 4

Hak Cipta sebagaimana dimaksud dalam Pasal 3 huruf a merupakan hak eksklusif yang terdiri atas hak moral dan hak ekonomi.

Fungsi dan sifat hak cipta Pasal 4

Ketentuan sebagaimana dimaksud dalam Pasal 23, Pasal 24, dan Pasal 25 tidak berlaku terhadap:

- a. Penggunaan kutipan singkat Ciptaan dan/atau produk Hak Terkait untuk pelaporan peristiwa aktual yang ditujukan hanya untuk keperluan penyediaan informasi aktual;
- b. Penggandaan Ciptaan dan/atau produk Hak Terkait hanya untuk kepentingan penelitian ilmu pengetahuan;
- c. Penggandaan Ciptaan dan/atau produk Hak Terkait hanya untuk keperluan pengajaran, kecuali pertunjukan dan Fonogram yang telah dilakukan Pengumuman sebagai bahan ajar; dan
- d. Penggunaan untuk kepentingan pendidikan dan pengembangan ilmu pengetahuan yang memungkinkan suatu Ciptaan dan/atau produk Hak Terkait dapat digunakan tanpa izin Pelaku Pertunjukan, Produser Fonogram, atau Lembaga Penyiaran.

Sanksi Pelanggaran Pasal 113

1. Setiap Orang yang dengan tanpa hak melakukan pelanggaran hak ekonomi sebagaimana dimaksud dalam Pasal 9 ayat (1) huruf i untuk Penggunaan Secara Komersial dipidana dengan pidana penjara paling lama 1 (satu) tahun dan/atau pidana denda paling banyak Rp100.000.000 (seratus juta rupiah).
2. Setiap Orang yang dengan tanpa hak dan/atau tanpa izin Pencipta atau pemegang Hak Cipta melakukan pelanggaran hak ekonomi Pencipta sebagaimana dimaksud dalam Pasal 9 ayat (1) huruf c, huruf d, huruf f, dan/atau huruf h untuk Penggunaan Secara Komersial dipidana dengan pidana penjara paling lama 3 (tiga) tahun dan/atau pidana denda paling banyak Rp500.000.000,00 (lima ratus juta rupiah).

**Rina Fitriana
Dedy Sugiarto
Miwan Kurniawan Hidayat
Yusri Eli Hotman Turnip**

Perancangan Model Sistem Intelijensia Bisnis



Perancangan Model Sistem Intelijensia Bisnis

Edisi Pertama
Copyright©2024
Cetakan Pertama: September, 2024

Ukuran: 15,5 cm x 23 cm; Halaman: xvi + 99

wi.2024.0490

Penulis:

- **Rina Fitriana**
- **Dedy Sugiarto**
- **Miwan Kurniawan Hidayat**
- **Yusri Eli Hotman Turnip**

Editor : Wahyu Kurniawadi
Cover : Maulana Arifin
Tata letak : Dita Yuni Setiawati

Penerbit
Wawasan Ilmu

Anggota IKAPI (215/JTE/2021)
Leler RT 002 RW 006 Desa Kaliwedi Kec. Kebasen Kab. Banyumas Jawa Tengah
53172
Email : redaksi@wawasanilmu.co.id
Web : <https://wawasanilmu.co.id/>

ISBN :

All Right Reserved

Hak Cipta pada Penulis
Hak Cipta Dilindungi Undang-undang

Dilarang memperbanyak sebagian atau seluruh isi buku ini dalam bentuk apapun, baik secara elektronis maupun mekanis, termasuk memfotokopi, merekam atau dengan sistem penyimpanan lainnya, tanpa izin tertulis dari penerbit.

PRAKATA

Kelebihan buku ini dibandingkan dengan buku lain adalah dilengkapi dengan contoh penulisan sebelumnya di industri manufaktur dan jasa di Indonesia. penulisan mengenai sistem intelijensia bisnis dan data analitik telah banyak dilakukan oleh Penulis. Sasaran utama pengguna buku ini adalah Dosen dan mahasiswa yang mengambil mata kuliah intelijensia bisnis dan Dosen dan mahasiswa yang mengambil topik perancangan model sistem intelijensia bisnis. Sasaran umum pengguna buku ini adalah pembaca yang tertarik dengan topik perancangan model sistem intelijensia bisnis.

Prasyarat pengguna buku ini adalah pembaca yang tertarik dengan topik sistem intelijensia bisnis dan berminat untuk penelitian dengan topik sistem intelijensia bisnis. Buku lain yang sudah dapat dijadikan pendamping agar buku lebih mudah dipahami adalah Data Mining dan Aplikasinya.

Ucapan terimakasih diberikan kepada semua pihak yang telah memberikan Hibah Buku Buku.

Jakarta, 1 Mei 2024

Penulis

Dummy Book Finish Wawasan Ilmu

DAFTAR ISI

PRAKATA.....	v
DAFTAR ISI.....	vii
DAFTAR GAMBAR	xi
DAFTAR TABEL	xv
BAB I	
PENDAHULUAN	1
BAB II	
SISTEM INTELIJENSIA BISNIS.....	5
A. Data Warehouse.....	8
B. Extract Transform Load.....	9
C. OLAP (Online Analytical Processing)	10
D. Cube.....	11
E. Business Analytics	13
F. <i>Data Mining</i>	14
1. K-Means Clustering	19

2. X-Means Clustering.....	21
G. Validasi dan verifikasi Model Sistem Intelijensia Bisnis.....	24
H. Validasi Clustering	26

BAB III

MODEL SISTEM INTELIJENSIA BISNIS31

A. Model Sistem Intelijensia Bisnis di Pendidikan Tinggi	31
1. Analisis Kebutuhan Sistem BI Publikasi Ilmiah Dosen.....	33
2. Model Sistem BI Publikasi Ilmiah Dosen.....	37
B. Studi Kasus di Industri Provider	46
1. Model Sistem Intelijensia Bisnis Kecepatan Provider	48
2. Proses ETL	49
3. Proses ETL untuk Tabel Dimensi	49
4. Proses Pembuatan Tabel Fakta	51
5. Perancangan Data Warehouse	52

BAB IV

PERANCANGAN SISTEM INTELIJENSIA BISNIS61

A. Perancangan Sistem Intelijensia Bisnis di Pendidikan Tinggi	61
1. Penerapan Clustering.....	62
2. Penerapan OLAP	65
3. Visualisasi Data.....	68
4. Evaluasi Sistem	72

B. Perancangan Sistem Intelijensia Bisnis di Industri
Provider79

1. Proses *Clustering*79

2. Penerapan OLAP82

3. Visualisasi Data.....83

DAFTAR PUSTAKA.....93

PROFIL PENULIS97

Dummy Book Finish Wawasan Ilmu

Dummy Book Finish Wawasan Ilmu

DAFTAR GAMBAR

Gambar 1.	Manfaat Sistem Intelijensia Bisnis (Vercellis, 2009).....	7
Gambar 2.	Contoh Model <i>Star Schema</i>	8
Gambar 3.	Contoh Model <i>Snowflake Schema</i>	9
Gambar 4.	Proses ETL Menggunakan Pentaho	10
Gambar 5.	OLAP Menggunakan Pentaho BI Server	11
Gambar 6.	Cube Modal Tertinggi	12
Gambar 7.	Operasi OLAP	12
Gambar 8.	Tiga Tipe Business Analytics (Sharda et al., 2018).....	14
Gambar 9.	Taksonomi Sederhana Pekerjaan Data Mining (Sharda et al., 2018).....	19
Gambar 10.	Ilustrasi Langkah-Langkah K-Means Clustering (Sharda et al., 2018).....	20
Gambar 11.	Ilustrasi Langkah-Langkah X-Means Clustering ..	22
Gambar 12.	Peta Clustering Pembagian Kabupaten	24
Gambar 13.	Model Resiko Kualitas Susu.....	25

Gambar 14. Model Sistem BI Publikasi Ilmiah Dosen	38
Gambar 15. Model Arsitektur Sistem BI Publikasi Ilmiah Dosen	39
Gambar 16. Model Arsitektur Data Warehouse Publikasi Ilmiah Dosen.....	39
Gambar 17. Model Dimensional Score Peneliti Menggunakan Star Schema	42
Gambar 18. Model Dimensional Score Peneliti Menggunakan Snowflakes Schema.....	42
Gambar 19. Komponen Step pada ETL fact_artikel.....	43
Gambar 20. Model Arsitektur BI.....	48
Gambar 21. Sistem BI.....	49
Gambar 22. Proses ETL pada Tabel dim_kecamatan.....	50
Gambar 23. Proses ETL pada Tabel dim_kelurahan	51
Gambar 24. Proses Pembentukan Tabel Fakta	52
Gambar 25. Proses Penentuan Primary Key pada tabel dimensi	54
Gambar 26. Proses penentuan foreign key	54
Gambar 27. <i>Star schema</i> untuk <i>data warehouse</i> data kecepatan <i>upload</i>	55
Gambar 28. <i>Star schema</i> untuk <i>data warehouse</i> data kecepatan <i>download</i>	55
Gambar 29. Proses penentuan primary key pada tabel dimensi	57
Gambar 30. Proses penentuan <i>foreign key</i>	58
Gambar 31. <i>Star schema</i> untuk <i>data warehouse</i> data kecepatan upload.....	58

Gambar 32. Star schema untuk data warehouse data kecepatan download	59
Gambar 33. Alur Proses Clustering	63
Gambar 34. Visualisasi Pengelompokan Menggunakan X-Means.....	64
Gambar 35. Cube Indeks Peneliti.....	66
Gambar 36. Cube Score Peneliti	66
Gambar 37. Cube Artikel Publikasi	67
Gambar 38. Cube Subjek Penelitian.....	67
Gambar 39. Tampilan Halaman Utama Sistem Intelijensia Bisnis.....	68
Gambar 40. Tampilan Halaman Indeks Peneliti	69
Gambar 41. Tampilan Halaman Score Peneliti	70
Gambar 42. Tampilan Halaman Artikel Publikasi	71
Gambar 43. Tampilan Halaman Subjek Penelitian.....	72
Gambar 44. Grafik Peringkat Persentil Nilai SUS	76
Gambar 45. Proses clustering pada data upload dan data download	80
Gambar 46. Titik centroid data upload	80
Gambar 47. Titik centroid data download	81
Gambar 48. OLAP Cube data kecepatan internet	82
Gambar 49. Tampilan dashboard awal atau overview	83
Gambar 50. Tampilan dashboard kecepatan upload	84
Gambar 51. Tampilan dashboard kecepatan download.....	86
Gambar 52. Tampilan dashboard kecepatan upload	86
Gambar 53. Tampilan dashboard kecepatan download.....	87

Gambar 54. Tampilan dashboard hasil clustering data
kecepatan upload..... 88

Gambar 55. Tampilan dashboard hasil clustering data
kecepatan download 88

Dummy Book Finish Wawasan Ilmu

DAFTAR TABEL

Tabel 1.	Support 1-Itemset.....	15
Tabel 2.	Support 1-Itemset Final	16
Tabel 3.	Total Pembentukan 2-Itemset.....	16
Tabel 4.	Support 2-Itemset Final	17
Tabel 5.	Pembentukan Aturan Asosiasi	17
Tabel 6.	Tabel jarak titik centroid antara cluster 1 dan cluster 2.....	23
Tabel 7.	Contoh Validasi Model Sistem Intelijensia Bisnis (Fitriana,2013)	25
Tabel 8.	Matriks Hasil Analisis PIECES.....	34
Tabel 9.	Relasi Tingkat Detail (Grain) dan Dimensi	40
Tabel 10.	Database Model Star Schema Database: publikasi_ star	44
Tabel 11.	Database Model Snowflakes Schema Database: publikasi_snow.....	44
Tabel 12.	Waktu Proses Query Model Star Schema	45
Tabel 13.	Waktu Proses Query Model Snowflakes Schema.....	45

Tabel 17.	Tabel grain dalam pembentukan data warehouse ...	56
Tabel 18.	Perbandingan Nilai Davies Bouldin Index.....	63
Tabel 19.	Hasil Clustering dan Centroid	64
Tabel 20.	Tingkat Pencapaian H-Index	65
Tabel 21.	Hasil Uji Verifikasi Sistem.....	73
Tabel 22.	Daftar Pertanyaan Kuesioner SUS	75
Tabel 23.	Jawaban Kuesioner SUS	77
Tabel 24.	Perhitungan Nilai SUS.....	77
Tabel 25.	Tabel informasi kecepatan upload tertinggi dan terendah.....	84
Tabel 26.	Tabel informasi kecepatan download tertinggi dan terendah	85
Tabel 27.	Informasi pakar uji validitas.....	89
Tabel 28.	Hasil Uji Validitas dashboard.....	90

BAB I

PENDAHULUAN

Kini semakin penting bagi dunia usaha dan dunia pendidikan untuk memiliki pandangan yang jelas atas semua data mereka agar tetap kompetitif, dan di situlah peran alat *Business Intelligence* (BI). Sudah banyak bisnis yang sudah menggunakan alat BI, dan proyeksi menunjukkan pertumbuhan yang berkelanjutan di bidang bisnis. tahun-tahun yang akan datang.

Namun bagi mereka yang belum menggunakan alat ini, atau hanya ingin mempelajari lebih lanjut, mungkin sulit untuk memahami secara pasti apa itu BI, dengan membaca buku buku ini diharapkan dapat memberi gambaran mengenai penerapan *Business intelligence*.

Business Intelligence (BI) atau Intelijensia Bisnis menurut Vercellis (2009) didefinisikan sebagai model matematika dan analisa metodologi eksploitasi data yang tersedia untuk mendapatkan informasi dan pengetahuan yang berguna dalam proses pengambilan keputusan yang kompleks. Intelijensia Bisnis adalah kombinasi dari data warehouse dan sistem intelijen. Sistem intelijensia bisnis dapat berperan serta sebagai alat untuk memberikan informasi yang akurat dan berguna bagi pengambil keputusan dalam batas waktu yang ditentukan untuk mendukung pengambilan keputusan.

Intelijensia bisnis menggabungkan analisis bisnis, penambangan data, visualisasi data, alat dan infrastruktur data, serta praktik terbaik untuk membantu organisasi membuat keputusan yang lebih berdasarkan data. Intelijensia bisnis modern memiliki pandangan komprehensif tentang data organisasi dan menggunakan data tersebut untuk mendorong perubahan, menghilangkan inefisiensi, dan dengan cepat beradaptasi terhadap perubahan pasar atau pasokan. Solusi BI modern memprioritaskan analisis layanan mandiri yang fleksibel, data yang diatur pada platform tepercaya, pemberdayaan pengguna bisnis, dan kecepatan dalam memperoleh wawasan.

Penelitian di bidang Sistem Intelijensia Bisnis telah banyak dilakukan oleh para peneliti diantaranya penelitian mengenai studi eksperimental melibatkan berbagai kumpulan data analisis keamanan intelijen bisnis, menilai metrik kinerja (Zhao L, Zhang J, 2024). Penelitian hubungan antara intelijensia bisnis, pembelajaran organisasi, jaringan pembelajaran, antisipasi nilai pelanggan, inovasi dan ekonomi kreatif berdasarkan performansi UMKM di Jawa Timur Indonesia (Anjaningrum *et.al*, 2024), Penelitian mengenai kerangka kerja untuk mengembangkan model kematangan spesifik domain, yang memenuhi industri perawatan kesehatan (Ramalingam *et. al.*, 2024)

Pendekatan manajerial diharapkan dengan adanya Intelijensia Bisnis dapat menghasilkan proses pengambilan keputusan yang efektif dan efisien bagi manajemen. Pendekatan teknis membahas peralatan dan software yang mendukung proses intelijensia bisnis. Pendekatan sistem membahas nilai tambah untuk mendukung sistem informasi. Beberapa Sistem Intelijensia Bisnis telah berhasil, salah satu faktornya mereka menggunakan software yang telah disediakan oleh Microsoft, Oracle, SAP dan lain-lain. Topik yang terintegrasi dengan Sistem Intelijensia Bisnis adalah Manajemen Rantai Pasok, Customer Relationship Management, Data Mining, Data Warehouse, Sistem Pendukung Keputusan, Performance Scorecard, Manajemen Pengetahuan, Business Process Management, Artificial Intelligence, Enterprise Resource Planning, Extract Transformation Loading, OLAP (On Line Analytical Processing), Sistem Manajemen Kualitas dan Manajemen Strategi. (Fitriana *et.al*, 2011)

Bisnis bisa berfokus pada keuntungan dan penjualan, tetapi intelijen bisnis berfokus pada data. Aktivitas yang bergantung pada data membutuhkan data analisis untuk dipelajari dari sumber yang berbeda. Buku ini membahas bagaimana mengerti big data dalam analisis intelijen bisnis, ekstrak data dari berbagai sumber, transforming data sehingga dapat dianalisis mengikuti kebutuhannya, loading data ke dalam sistem intelijen bisnis untuk dianalisis. Buku Perancangan Model Sistem Intelijen Bisnis ini ditulis berdasarkan penelitian yang telah dilakukan sebelumnya, sehingga diharapkan dapat menjadi contoh untuk pengembangan penelitian sistem intelijen bisnis. Contoh-contoh studi kasus analisis kebutuhan, metode, model dan sistem intelijen bisnis dibahas di buku ini.

Dummy Book Finish Wawasan Ilmu

BAB II

SISTEM INTELIJENSIA BISNIS

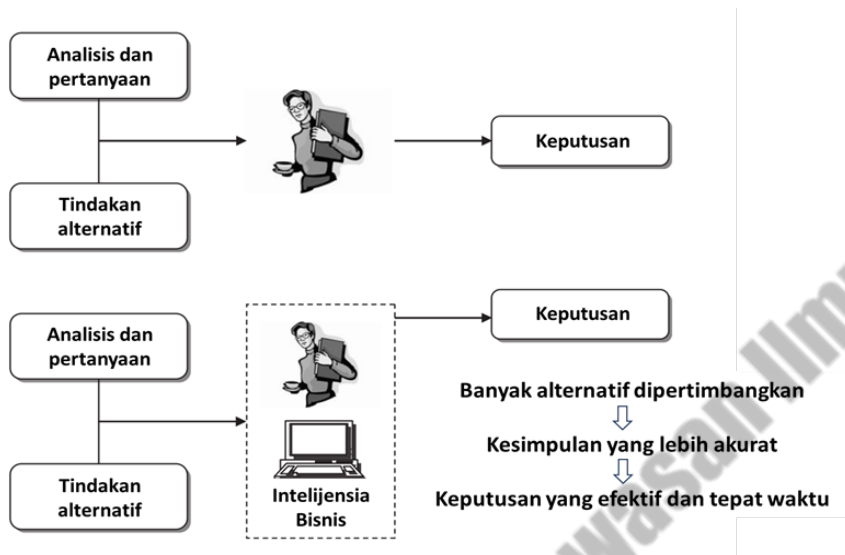
Business Intelligence (BI) atau Intelijensia Bisnis menurut Vercellis (2009) didefinisikan sebagai model matematika dan analisa metodologi eksploitasi data yang tersedia untuk mendapatkan informasi dan pengetahuan yang berguna dalam proses pengambilan keputusan yang kompleks. Intelijensia Bisnis adalah kombinasi dari *data warehouse* dan sistem intelijen. Sistem intelijensia bisnis dapat berperan serta sebagai alat untuk memberikan informasi yang akurat dan berguna bagi pengambil keputusan dalam batas waktu yang ditentukan untuk mendukung pengambilan keputusan dalam Perusahaan.

Sistem Intelijensia Bisnis adalah berguna bagi perusahaan, dimana para pemimpin dan manajer dapat mengambil keputusan yang lebih baik berdasarkan data yang akurat. Sistem intelijensia membantu perusahaan untuk bekerja lebih efisien dengan mengidentifikasi cara-cara untuk menghemat waktu dan sumber daya. Dengan memahami pasar dan persaingan lebih baik, perusahaan dapat segera menyesuaikan diri terhadap perubahan. Selain itu, alat ini memudahkan bagi tim dalam organisasi untuk bekerja sama dan berbagi informasi. Sistem kecerdasan bisnis bisa membantu perusahaan meningkatkan kinerja mereka secara keseluruhan.

Intelijensia Bisnis adalah istilah yang mencakup basis data, struktur, program aplikasi, instrumen analisis, dan pendekatan yang bertujuan untuk memungkinkan akses data secara interaktif, memfasilitasi pengolahan data, sehingga para analis dan manajer bisnis dapat melakukan analisis dengan efektif (Sharda *et al.*, 2018). Sistem Intelijensia Bisnis mengintegrasikan data operasional dengan alat analitik untuk menyajikan informasi yang rumit dan bersaing kepada para perencana dan pengambil keputusan (Negash, 2004).

Manfaat dari penggunaan Intelijensia Bisnis mencakup kemampuan untuk melihat laporan secara independen; mengidentifikasi potensi pemborosan dalam sistem; mengenali kekuatan dan kelemahan perusahaan; meningkatkan efisiensi dalam pengambilan keputusan; memungkinkan analisis *real-time* dengan navigasi yang cepat; memudahkan berbagi dan mengakses informasi; memberikan respons cepat terhadap pertanyaan dan permasalahan bisnis; menyediakan visualisasi data agar mudah dibaca, dimengerti, dan ditafsirkan (Joshi & Dubbewar, 2021).

Gambar 1 menunjukkan manfaat utama yang dapat diperoleh oleh suatu organisasi melalui pengadopsian Sistem Intelijensia Bisnis. Ketika menghadapi masalah, para pengambil keputusan mengajukan serangkaian pertanyaan dan melakukan analisis yang sesuai. Oleh karena itu, mereka meneliti dan membandingkan beberapa opsi, memilih keputusan terbaik, dengan mempertimbangkan kondisi saat itu. Jika para pengambil keputusan dapat mengandalkan Sistem Intelijensia Bisnis yang mempermudah aktivitas mereka, kita dapat mengharapkan bahwa kualitas keseluruhan proses pengambilan keputusan akan meningkat secara signifikan. Dengan bantuan model matematika dan algoritma, sebenarnya memungkinkan untuk menganalisis sejumlah besar tindakan alternatif, mencapai kesimpulan yang lebih akurat, dan membuat keputusan yang efektif dan tepat waktu. Oleh karena itu, dapat disimpulkan bahwa keuntungan utama dari pengadopsian sistem kecerdasan bisnis terletak pada peningkatan efektivitas proses pengambilan keputusan.



Gambar 1. Manfaat Sistem Intelijensia Bisnis (Vercellis, 2009)

Intelijensia Bisnis menurut Niu (2009) merujuk kepada proses untuk mengekstrak, mentransformasi, memanajemen dan menganalisa data bisnis untuk mendukung pembuatan keputusan. Proses ini berdasarkan pada database yang besar yang dibentuk menjadi *data warehouse*, dengan misi untuk menyebarkan intelijen atau pengetahuan di dalam seluruh organisasi, dari level strategik terhadap level takstis dan operasional

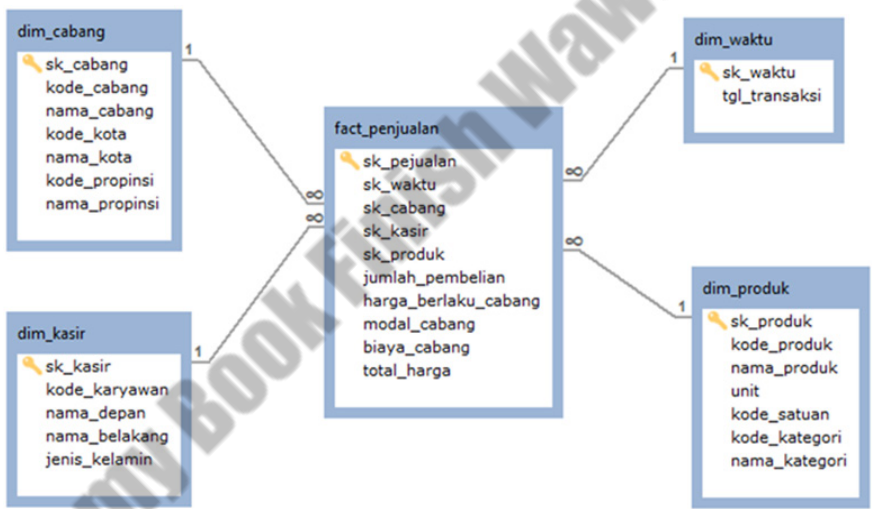
Sistem Intelijensia Bisnis (BI) menurut Niu et al. (2009) adalah data pendukung dari Sistem Penunjang Keputusan (SPK) yang berfokus pada manipulasi dari volume data perusahaan yang besar dalam data warehouses. Menurut Vercellis (2009) tujuan utama sistem Intelijensia Bisnis adalah memberikan alat dan metodologi kepada karyawan sehingga dapat membuat mereka dapat mengambil keputusan yang efektif dan efisien. Keuntungan Sistem Bisnis Intelijen yaitu lebih banyak alternatif dipertimbangkan, kesimpulan yang lebih akurat serta keputusan dan waktu lebih efektif.

Peran Sistem Intelijensia Bisnis adalah meningkatkan kepuasan pelanggan dan stakeholder, mendukung pengambilan keputusan

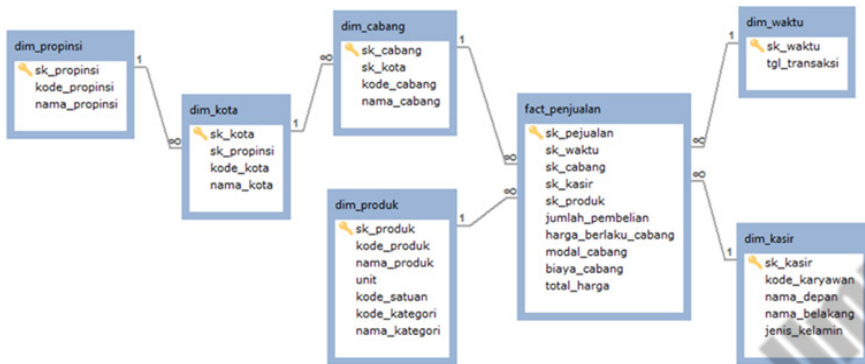
secara intelijen, mengembangkan alat-alat pendukung keputusan mengembangkan alat-alat pendukung Keputusan, membuat optimisasi dan mengembangkan model. (Fitriana *et.,al*,2012)

A. Data Warehouse

Data warehouse adalah tempat penyimpanan data yang menggabungkan informasi dari berbagai sumber dengan tujuan untuk melakukan analisis data berdimensi khusus. Rancangan *relational data warehouse* direpresentasikan dalam model dimensional yang terdiri dari *star schema* dan *snowflake schema*. Perbedaan antara model *star schema* dengan *snowflake schema* dicontohkan pada Gambar 2 dan Gambar 3.



Gambar 2. Contoh Model *Star Schema*

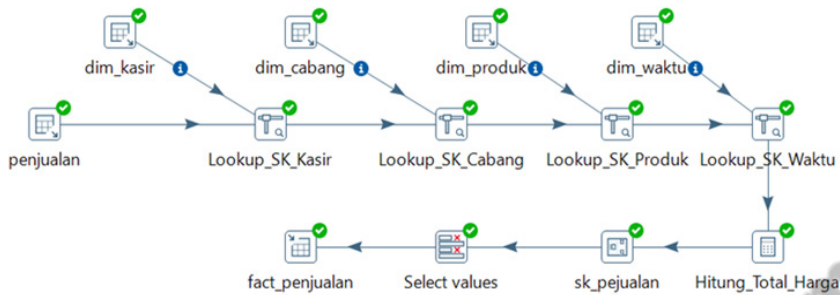


Gambar 3. Contoh Model Snowflake Schema

Star schema terdiri dari satu tabel fakta terpusat dan satu set tabel dimensi. Secara umum tabel dimensi pada *star schema* tidak dinormalisasi dan memungkinkan berisi *redundant data*. *Snowflake schema* menghindari redundansi pada *star schema* melalui representasi dimensi berupa beberapa tabel yang terkait dengan *referential integrity constraints* (Vaisman & Zimányi, 2014). Kinerja model dimensional yang lebih baik dapat ditentukan dari ukuran *data warehouse* yang lebih kecil dan waktu pemrosesan *query* yang lebih singkat. Ukuran *data warehouse* yang kecil mengarah pada konsumsi memori yang lebih sedikit (Iqbal et al., 2020).

B. Extract Transform Load

Proses ETL terdiri dari ekstraksi (membaca data dari satu atau lebih database), transformasi (mengubah data yang diekstraksi dari bentuk sebelumnya ke dalam bentuk yang diperlukan sehingga dapat ditempatkan ke dalam *data warehouse*), dan *load* (menempatkan data ke dalam *data warehouse*). Pada Gambar 4 adalah proses ETL menggunakan perangkat lunak Pentaho Data Integration (Pentaho).



Gambar 4. Proses ETL Menggunakan Pentaho

Data yang digunakan dalam proses ETL dapat berasal dari berbagai sumber, misal aplikasi ERP, alat CRM, *flat file*, Excel *spreadsheet*. Tujuan dari proses ETL adalah terbentuknya *data warehouse* dengan data terintegrasi dan bersih.

C. OLAP (Online Analytical Processing)

Data yang telah disimpan di *data warehouse* dapat digunakan dengan berbagai cara untuk mendukung pengambilan keputusan organisasi. OLAP bisa dikatakan teknik analisis data yang paling umum digunakan di *data warehouse*. Penerapan OLAP merupakan proses pembentukan *cube* OLAP sebagai representasi data pada model multidimensi untuk memenuhi kebutuhan informasi yang diperlukan dalam mendukung pengambilan keputusan. Gambar 5 merupakan tampilan OLAP menggunakan perangkat lunak Pentaho BI Server.

Cabang	Waktu	Produk	Measures • Jumlah Pembelian
▣ All Dim Cabang.Cabangs	▣ All Dim Waktu.Waktus	▣ All Dim Produk.Produks	4,601,777
PHI Mini Market - Jakarta Pusat 01	▣ All Dim Waktu.Waktus	▣ All Dim Produk.Produks	1,531,354
PHI Mini Market - Makassar 01	▣ All Dim Waktu.Waktus	▣ All Dim Produk.Produks	1,531,616
PHI Mini Market - Surabaya 01	▣ All Dim Waktu.Waktus	▣ All Dim Produk.Produks	1,538,807
	2008-01-01 00:00:00.0	▣ All Dim Produk.Produks	3,997
		Manggis 1 kg	89
		air mineral 600 ml	41
		alpukat 1 kg	88
		apel 1 kg	184
		bawang merah 1kg	63
		bawang putih 1 kg	89
		belimbing 1 kg	124
		bumbu fried chicken	57
		buncis 1 kg	146
		durian 1 kg	146
		jeruk 1 kg	58
		jus kesehatan 600 ml	118
		kacang goreng 1 kg	61
		kacang hijau 1 kg	75
		kacang mete 1 kg	99
		kacang panjang 1 kg	106
		kelapa muda	82
		kentang 1 kg	149
		ketimun 1 kg	98

Gambar 5. OLAP Menggunakan Pentaho BI Server

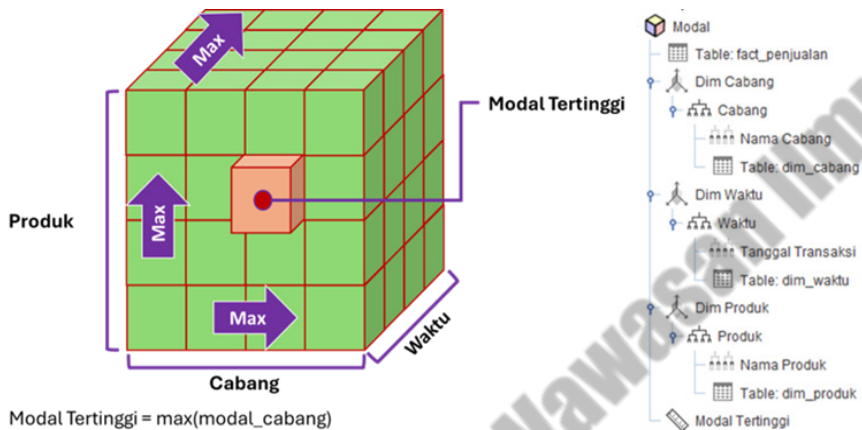
Alur kerja penerapan OLAP melalui empat langkah utama, yaitu:

- Mengekstraksi data dari sumber asli dan memuat data ke dalam *data warehouse*.
- Menyiapkan data (*integrating, cleaning, filtering, aggregating, dan joining*).
- Membentuk *cube* OLAP.
- Mengakses data untuk analisis.

D. Cube

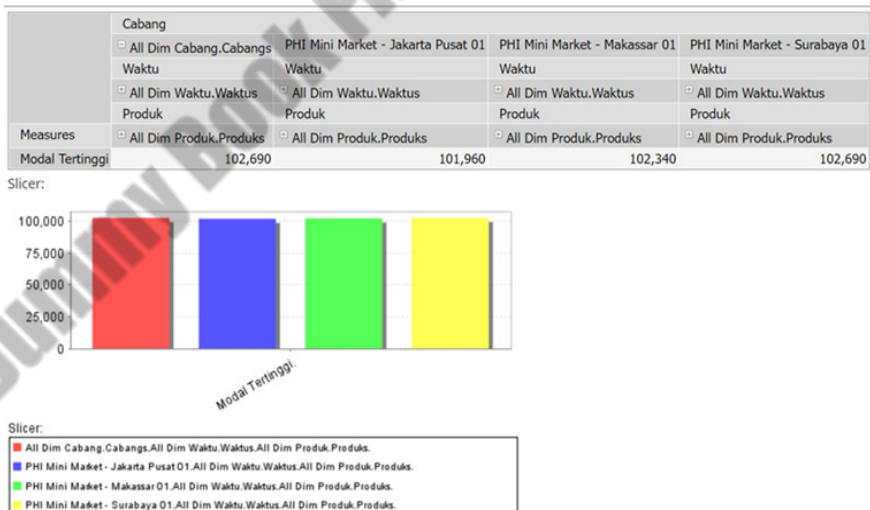
Penggunaan OLAP bagi seorang analis yaitu untuk melakukan navigasi database pada subset tertentu dari data dan perkembangannya dari waktu ke waktu dengan mengubah orientasi data dan mendefinisikan perhitungan analitis. Struktur operasional utama dalam OLAP didasarkan pada konsep yang

disebut *cube*. *Cube* dalam OLAP adalah struktur data multidimensi yang memungkinkan analisis data dengan cepat (Sharda et al., 2018). Contoh penggunaan *cube* yaitu untuk memberikan informasi tentang jumlah modal tertinggi dapat dilihat pada Gambar 6.



Gambar 6. Cube Modal Tertinggi

Operasi OLAP yang umum digunakan meliputi *slice and dice*, *roll-up*, *drill down* and *pivot*, seperti contoh pada Gambar 7.



Gambar 7. Operasi OLAP

E. Business Analytics

Terdapat tiga tipe *business analytics* yaitu *descriptive*, *predictive*, *prescriptive*.

1. *Descriptive Analytics*

Descriptive analytics berupa pelaporan yang merujuk pada pengetahuan dan pemahaman dari sesuatu yang terjadi disertai dasar penyebabnya dalam organisasi. Tujuan dari tugas deskriptif adalah untuk menurunkan pola-pola (korelasi, *trend*, *cluster*, trayektori, dan anomali) yang meringkas hubungan yang pokok dalam data. Tugas *data mining* deskriptif sering merupakan penyelidikan dan seringkali memerlukan teknik *postprocessing* untuk validasi dan penjelasan hasil. (Fitriana et.,al.,2022)

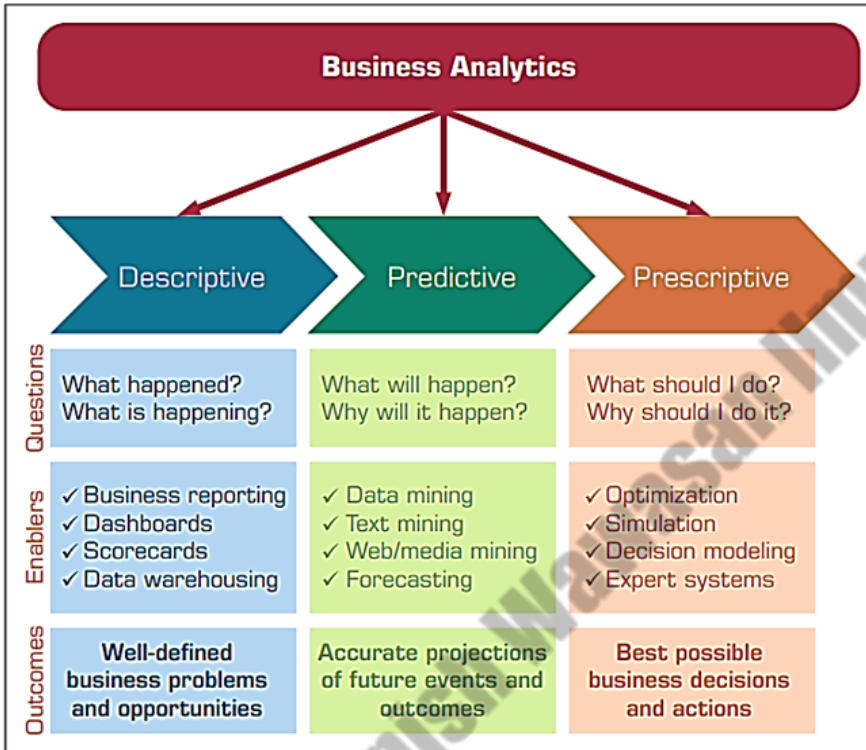
2. *Predictive Analytics*

Predictive analytics memiliki tujuan untuk mengenali kemungkinan yang akan terjadi pada masa selanjutnya. *Predictive analytics* berdasarkan pada cara statistik dan cara lain seperti *data mining*.

3. *Prescriptive Analytics*

Tujuan analitik preskriptif yaitu mengidentifikasi sesuatu yang sedang terjadi dan menciptakan keputusan untuk memperoleh hasil kerja yang paling baik.

Ketiga tipe *business analytics* diilustrasikan pada Gambar 8



Gambar 8. Tiga Tipe Business Analytics (Sharda et al., 2018)

F. *Data Mining*

Data mining adalah suatu proses untuk mendapatkan informasi yang bernilai dalam repositori data yang berukuran besar secara otomatis (Tan et al., 2019). Secara umum *data mining* dibagi menjadi tiga kategori, yaitu:

1. *Prediction*

Prediction biasa disebut sebagai tindakan menceritakan tentang masa depan dengan mempertimbangkan pengalaman, pendapat, dan informasi lain yang relevan dalam melakukan tugas ramalan.

2. Association

Asosiasi atau pembelajaran aturan asosiasi dalam penambangan data, adalah teknik yang populer dan diteliti dengan baik untuk menemukan hubungan yang menarik di antara variabel dalam *database* besar.

Contoh penelitian Asosiasi menganalisis transaksi dari transaksi kedai kopi dalam enam bulan untuk mengetahui pola pembelian konsumen dengan pengembangan model aturan asosiasi. Salah satu kedai kopi di Jakarta yaitu 8th Bean Cafe dijadikan studi kasus dalam penelitian ini. Permasalahan yang terjadi saat pengambilan keputusan tidak berdasarkan analisa data sehingga hilang selama 3 periode penjualan. Kafe ini memiliki 80 menu yang terus berubah setiap saat sesuai keinginan pemiliknya tanpa mengetahui menu favorit, menu yang paling sering dibeli, dan lain sebagainya. Mereka tidak pernah menganalisis menu apa yang berminat dibeli konsumen.

Tabel 1. Support 1-Itemset

Itemset	Support
N	0.333
A	0.833
J	0.667
B	0.333
K	0.667
O	0.167

$$* \text{Support}(N) = \frac{\text{Jumlah transaksi yang mengandung } N}{\text{Total transaksi}} = \frac{2}{6} = 0.333$$

$$* \text{Support}(A) = \frac{\text{Jumlah transaksi yang mengandung } A}{\text{Total transaksi}} = \frac{5}{6} = 0.833$$

Hasil nilai *support* pada calon 1-itemset ditunjukkan oleh tabel 1. Itemset O memiliki nilai *support* di bawah minimum *support* sebesar 0.167. Sehingga itemset O harus dihilangkan dan tidak dilanjutkan

ke perhitungan berikutnya.

Tabel 2. Support 1-Itemset Final

Itemset	Support
N	0.333
A	0.833
J	0.667
B	0.333
K	0.667

$$F1 = \{\{N\}, \{A\}, \{J\}, \{B\}, \{K\}\}$$

Tabel 3. Total Pembentukan 2-Itemset

Itemset	Total
N, A	1
N, J	1
N, B	1
N, K	1
A, J	3
A, B	2
A, K	3
J, B	0
J, K	3
B, K	1

Tabel 4. Support 2-Itemset Final

Itemset	Support
A, J	0.500
A, B	0.333
A, K	0.500
J, K	0.500

$$F2 = \{\{A,J\}, \{A,B\}, \{A,K\}, \{J,K\}\}$$

Adapun contoh perhitungan *support* dari masing-masing kombinasi itemset adalah sebagai berikut:

$$\begin{aligned}
 * \text{Support } (A, J) &= \frac{\sum \text{Jumlah transaksi mengandung } A \text{ dan } J}{\text{Total transaksi}} \\
 &= \frac{3}{6} = 0.500
 \end{aligned}$$

Tabel 5. Support 3-Itemset Final

Itemset	Support
A, J, K	0.333

Frequent 3-itemset yang terbentuk (F3) yaitu {A,J,K}.

Tabel 5. Pembentukan Aturan Asosiasi

Aturan	Support	Confidence
Jika membeli A, maka akan membeli J	0.500	0.600
Jika membeli J, maka akan membeli A	0.500	0.750
Jika membeli A, maka akan membeli B	0.333	0.400
Jika membeli B, maka akan membeli A	0.333	1.000
Jika membeli A, maka akan membeli K	0.500	0.600
Jika membeli K, maka akan membeli A	0.500	0.750

Jika membeli J, maka akan membeli K	0.500	0.750
Jika membeli K, maka akan membeli J	0.500	0.750
Jika membeli A dan J, maka akan membeli K	0.333	0.667

Berdasarkan hasil Asosiasi menggunakan algoritma apriori menghasilkan 3 aturan asosiasi final. Aturan asosiasi yang terbentuk diantaranya $\{B \Rightarrow A\}$, $\{J \Rightarrow K\}$, $\{K \Rightarrow J\}$ atau {CLASSIC SIGNATURE}, {FRIED RICE AND PASTA LIGHT BITES}, {LIGHT BITES FRIED RICE AND PASTA} dari data mining transaksi penjualan. Aturan tersebut memberikan informasi bahwa dua jenis kombinasi 2 itemset cenderung dibeli secara bersamaan. (Elfira et.al., 2021)

3. *Clustering*

Clustering mempartisi kumpulan hal (misalnya, objek, peristiwa, disajikan dalam kumpulan data terstruktur) menjadi segmen (pengelompokan alami) yang anggotanya memiliki karakteristik yang sama. Berbeda dengan klasifikasi, dalam pengelompokan label kelas tidak diketahui.

Gambar 9 menunjukkan taksonomi sederhana untuk pekerjaan *data mining*, bersama dengan metode pembelajaran dan algoritma populer.

Data Mining Tasks & Methods	Data Mining Algorithms	Learning Type
Prediction		
Classification	Decision Trees, Neural Networks, Support Vector Machines, kNN, Naive Bayes, GA	Supervised
Regression	Linear/Nonlinear Regression, ANN, Regression Trees, SVM, kNN, GA	Supervised
Time series	Autoregressive Methods, Averaging Methods, Exponential Smoothing, ARIMA	Supervised
Association		
Market-basket	Apriori, OneR, ZeroR, Eclat, GA	Unsupervised
Link analysis	Expectation Maximization, Apriori Algorithm, Graph-Based Matching	Unsupervised
Sequence analysis	Apriori Algorithm, FP-Growth, Graph-Based Matching	Unsupervised
Segmentation		
Clustering	k-means, Expectation Maximization (EM)	Unsupervised
Outlier analysis	k-means, Expectation Maximization (EM)	Unsupervised

Gambar 9. Taksonomi Sederhana Pekerjaan Data Mining (Sharda et al., 2018)

1. K-Means Clustering

K-Means clustering adalah algoritma *clustering* yang dimulai dengan menetapkan jumlah *K* sebagai jumlah *cluster*. Selanjutnya menentukan nilai awal *centroid* dan setiap data diarahkan ke *centroid* terdekat. Setelah *cluster* terbentuk, *centroid* untuk setiap *cluster* diperbarui (Yusuf et al., 2022). Berikut ini adalah langkah-langkah algoritma pengelompokan *K-Means* (Sharda et al., 2018):

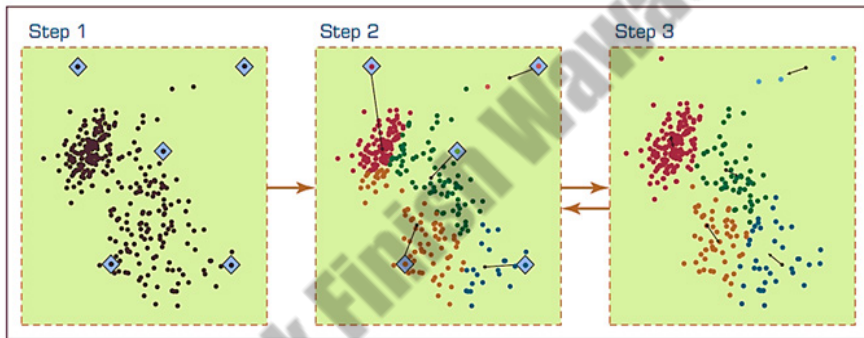
Langkah inisialisasi: Pilih jumlah *cluster* (yaitu nilai k).

Langkah 1: Secara acak menghasilkan k titik acak sebagai pusat *cluster* awal.

Langkah 2: Tetapkan setiap titik ke pusat *cluster* terdekat.

Langkah 3: Hitung ulang pusat *cluster* baru.

Langkah pengulangan: Ulangi langkah 2 dan 3 sampai beberapa kriteria konvergensi terpenuhi (perulangan berhenti jika nilai *centroid* yang dihasilkan tetap dan anggota *cluster* tidak berpindah ke *cluster* lain). Langkah-langkah algoritma pengelompokan *K-Means* diilustrasikan pada Gambar 10.



Gambar 10. Ilustrasi Langkah-Langkah K-Means Clustering (Sharda et al., 2018)

Centroid awal ditentukan dengan cara mengambil acak dari data yang tersedia. Selanjutnya pada tahap iterasi menggunakan rumus berikut (Putra & Kadyanan, 2021):

$$v_{ij} = \frac{1}{N_i} \sum_{k=0}^{N_i} x_{kj} \quad (2.1)$$

v_{ij} = *centroid* pada *cluster* ke- i untuk variabel ke- j

N_i = jumlah data yang merupakan anggota *cluster* ke- i

i, k = indeks *cluster*

j = indeks variabel

x_{ij} = nilai data ke- k yang ada di dalam *cluster* tersebut untuk variabel ke- j

Jarak antara setiap titik objek dengan *centroid* dihitung untuk menemukan jarak terdekat, pada umumnya dihitung menggunakan persamaan *Euclidean Distance* yaitu (Yusuf et al., 2022):

$$d_{(i,j)} = \sqrt{(x_{i1} - x_{j1})^2 + (x_{i2} - x_{j2})^2 + \dots + (x_{ki} - x_{kj})^2} \quad (2.2)$$

$d_{(i,j)}$ = jarak data ke- i ke pusat *cluster*- j

x_{ki} = data ke- i pada atribut data ke- k

x_{kj} = titik pusat ke- j pada atribut ke- k

2. X-Means Clustering

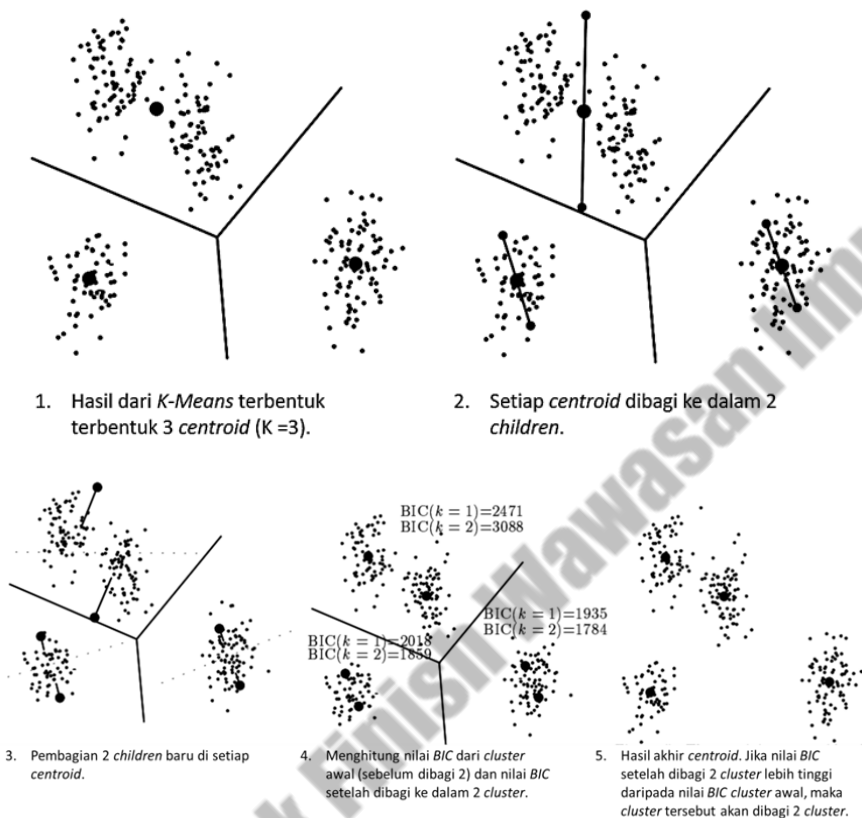
Algoritma pengelompokan *X-Means* adalah pengelompokan *K-Means* yang diperluas. *Bayesian Information Criterion (BIC)* digunakan pada algoritma *X-Means* untuk mengenali jumlah *cluster* yang tepat secara otomatis. Penentuan K *cluster* pada pengelompokan *X-Means* lebih dinamis dibandingkan dengan *K-Means*. Pada tahap awal, harus menetapkan nilai minimal $K_{\min}$ serta nilai maksimal $K_{\max}$. Nilai K mana yang harus dipilih dalam rentang $[K_{\min}, K_{\max}]$ akan teridentifikasi oleh *X-Means*. Tujuan dari algoritma *X-Means* adalah untuk menentukan nilai K dan waktu proses secara efisien pada jumlah data yang besar. Secara garis besar *X-Means* mencakup dua langkah. Langkah 1, disebut *Improve-Params*, menjalankan *K-Means* hingga konvergen. Langkah 2, disebut *Improve-Structure*, memutuskan apakah sebuah *cluster* harus dipecah menjadi dua *sub-cluster* atau tidak berdasarkan *BIC* (Yusuf et al., 2022). Berikut ini adalah langkah-langkah algoritma pengelompokan *X-Means*:

Langkah 1: *Improve-Params*

Langkah 2: *Improve-Structure*

Jika $K > K_{\max}$, mengembalikan skor terbaik. Jika tidak, lanjutkan ke langkah 1.

Langkah-langkah algoritma pengelompokan *X-Means* diilustrasikan pada Gambar 11.



Gambar 11. Ilustrasi Langkah-Langkah X-Means Clustering (Wijayanto & Adhitama, 2019)

BIC score dihitung berdasarkan persamaan dari Kass dan Wasserman sebagai berikut (Wijayanto & Adhitama, 2019):

$$BIC(M_j) = \hat{l}_j(D) - \frac{P_j}{2} \cdot \log R \quad (2.3)$$

$\hat{l}_j(D)$ merupakan fungsi *log-likelihood* data berdasarkan model ke- j , P_j adalah jumlah parameter pada M_j dan R diganti dengan total jumlah titik yang termasuk dalam *centroid* yang ditinjau.

Contoh kasus pembagian wilayah kabupaten berdasarkan luas area hutan antara lain hutan lindung, kawasan perlindungan,

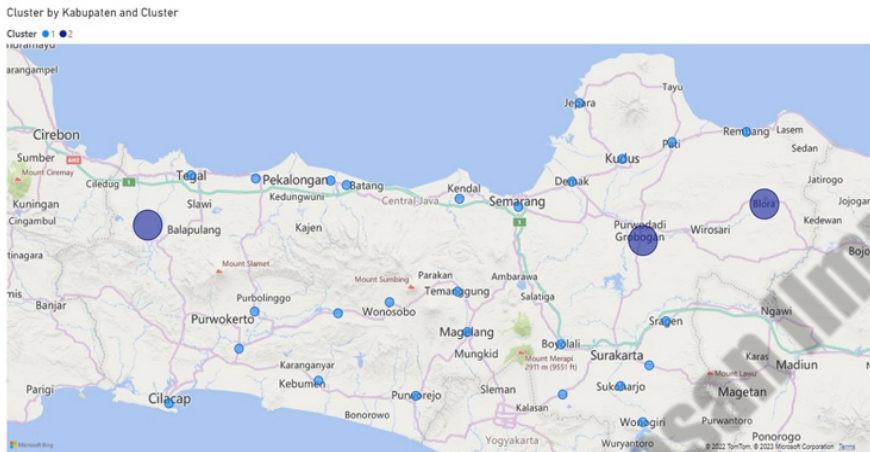
kawasan untuk produksi, dan kawasan untuk pengguna lain di provinsi Jawa Tengah menggunakan metode *Data Mining* yaitu *K-Means*. (Yusri & Fitriana, 2021). Data diperoleh dari Open data Jabar untuk area Jawa Tengah dimana terdapat empat jenis hutan yang akan dikelompokkan. Berdasarkan indeks Davies Bouldin terkecil, yaitu sebesar 0,436 pada pengelompokan dengan 2 cluster untuk pengelompokan kabupaten berdasarkan luar hutan. Kedua *cluster* dibedakan berdasarkan nilai jarak kedekatan *attribute* jenis hutan dengan titik *centroid* pada masing-masing *cluster*. Tabel 1 menunjukkan jarak titik centroid antara cluster 1 dan 2.

Tabel 6. Tabel jarak titik centroid antara cluster 1 dan cluster 2

Attribute	Cluster 1	Cluster 2
Hutan Lindung	3556,054	4521,584
Kawasan Perlindungan	3436,483	7075,780
Kawasan untuk Produksi	9098,848	54679,723
Kawasan untuk Pengguna Lain	725,705	3283,207

Proses *clustering* yang dilakukan mengelompokkan 26 kabupaten di wilayah provinsi Jawa Tengah ke dalam *cluster* 1, sedangkan *cluster* 2 terdiri dari 3 kabupaten di wilayah Jawa Tengah yaitu kabupaten Grobogan, kabupaten Blora, dan kabupaten Brebes. Visualisasi menggunakan *Power Business Intelligence*.

CLUSTERING KABUPATEN BERDASARKAN LUAS HUTAN DI PROVINSI JAWA TENGAH



Gambar 12. Peta Clustering Pembagian kabupaten

G. Validasi dan verifikasi Model Sistem Intelijensia Bisnis

Tes Verifikasi adalah metode pengetesan yang efektif untuk menghilangkan cacat di dalam *software*. Sedangkan tes validasi adalah sebuah metodologi untuk mengetes *software* (Perry, 2006).

Pengenalan data yang tersedia model matematika yang diadopsi dapat didefinisikan, dengan memastikan ketersediaan data yang dibutuhkan untuk setiap model dan verifikasi efisiensi algoritma yang akan digunakan akan cukup memadai untuk menyelesaikan besarnya persoalan.

Sistem Intelijensia Bisnis diverifikasi dengan jalan menguji apakah program untuk submodel tersebut telah berjalan dengan baik dan benar. Hal ini dilakukan dengan memberikan data input kepada setiap model program Sistem Intelijensia Bisnis dan hasil outputnya diperiksa apakah telah sesuai dengan hasil perhitungan memakai rumus. Bila masih ada kesalahan, maka program diperiksa dan diperbaiki. Setelah program berjalan dengan baik, maka akan ditentukan parameter-parameter submodel yang memberikan hasil paling optimal.

No	Fungsi Proses	Jenis Kegagalan	Efek Kegagalan Proses	Penyebab Kegagalan Proses	Severity	Occurance	Detectability	Submit	RPN	FRPN	Rekomendasi Tindakan
1	Uj Total Plate Control (TPC) agar diketahui TPC susu segar dari petani	Kandungan Total Plate Control (TPC) lebih besar dari 3 jutami	Tidak memenuhi standar mutu yang diminta IPS	High Sanitasi di tingkat peternak dari TPC belum sepenuhnya terlaksana dengan baik	7	7	7	Submit Query	343	692.1	Dibankan Standar Operating Proccdur di tingkat peternak. Sapi dimandikan, tangan pemerah dicuci sebelum pemerah susu, ember dibersihkan.
				Celup puting pasca pemerahan belum dilaksanakan oleh seluruh peternak	7	7	7	Submit Query	343	692.1	Dibuat Standar Operating Prosedur untuk Pemerahan Susu
				Adanya kotoran susu dari peternak melebihi 60 menit untuk sampai di cooling unit KPS	7	7	7	Submit Query	343	692.1	Dibuat Standar Operating Prosedur untuk Pengiriman Susu
		Kandungan kuman lebih besar dari 5 jutami susu	Dikenakan profit Rp. 100/kg	Kandang, Tangan manusia atau Ember kotor	6	6	6	Submit Query	216	593	Dibuat Standar Operating Prosedur untuk higienisasi tingkat peternak dan TPS
2	Uj Total Solid: agar diketahui jumlah Total Solid susu segar dari petani	Kandungan Total Solid kurang dari 11,3 %	Harga susu lebih rendah	Kurang konsentrat makanan	6	6	6	Submit Query	216	593	Sapi diberi makanan konsentrat, anggota koperasi mendapat subsidi makanan konsentrat
3	Uj Fat: agar diketahui jumlah Fat susu segar dari petani	Kandungan Fat kurang dari 3%	Harga susu lebih rendah	Kurang konsentrat makanan	3	3	3	Submit Query	27	307.9	Sapi diberi makanan konsentrat, anggota koperasi mendapat subsidi makanan konsentrat
4	Uj kandungan protein: agar diketahui kandungan protein susu segar dari petani	Kandungan Protein kurang dari 2,7% /ml susu	Harga susu lebih rendah	Kurang konsentrat makanan	5	3	5	Submit Query	75	500	Sapi diberi makanan konsentrat, anggota koperasi mendapat subsidi makanan konsentrat
5	Uj antibiotik: agar diketahui kandungan antibiotik susu segar dari petani	Susu mengandung antibiotik	susu ditolak	Sapi diben obat yang mengandung antibiotik	3	3	4	Submit Query	36	307.9	Susu yang mengandung antibiotik diberi tanda agar dipisahkan untuk diberikan ke anak sapi (pullet)
6	Uj Pemalsuan: agar diketahui adanya pemalsuan susu segar dari petani	Susu dicampur dengan air	Penurunan mutu susu	Pemalsuan susu dengan air	3	3	3	Submit Query	27	307.9	Peneguran kepada peternak dan penali harga

Gambar 13. Model Resiko Kualitas Susu

Validasi Model Sistem Intelijensia Bisnis menggunakan teknik *face validity* (Sargent, 1999; Sekaran,2000).

Tabel 7. Contoh Validasi Model Sistem Intelijensia Bisnis (Fitriana,2013)

No	Model	Pakar 1	Pakar 2
1.	Model Kualitas dengan metode <i>Fuzzy FMEA</i> (Failure Mode Effect Analysis)	Sangat setuju	Sangat Setuju

2.	Model CRM dengan metode RFM (Recency Frekuensi Monetary, CLV (Customer LIFE-TIME dan OLAP Cube.	Sangat setuju	Sangat Setuju
3.	Model Perkreditan dengan metode <i>Data Mining Decision Tree</i> .	Sangat Setuju	Sangat Setuju

H. Validasi Clustering

Validasi *clustering* penting untuk keberhasilan aplikasi pengelompokan. Ukuran numerik yang disebut juga sebagai kriteria atau indeks, diterapkan untuk menilai berbagai aspek validitas *cluster*. Validasi *clustering* dapat dikategorikan menjadi dua jenis utama (Kantardzic, 2020):

- *Internal measures*: Digunakan untuk mengukur struktur pengelompokan yang baik tanpa menggunakan informasi eksternal atau ketika informasi eksternal tidak tersedia. Langkah-langkah validasi internal hanya mengandalkan informasi dalam data.
- *External measures*: Digunakan untuk mengukur sejauh mana label *cluster* cocok dengan label kelas yang disediakan secara eksternal. Misalnya, sampel yang diberi label kelas sebagai “ground truth”.

Setidaknya ada 12 *internal measures* untuk validasi *clustering* (Xiong & Li, 2014), yaitu:

- *Root-mean-square standard deviation (RMSSTD)*, akar kuadrat dari varians sampel gabungan dari semua atribut untuk mengukur homogenitas *cluster* yang terbentuk.
- *R-squared (RS)*, rasio jumlah kuadrat antara *cluster* dengan jumlah total kuadrat dari seluruh kumpulan data untuk mengukur tingkat perbedaan antara *cluster*.
- *Modified Hubert Γ statistic (Γ)*, mengevaluasi perbedaan antara

cluster dengan menghitung ketidaksepakatan pasangan objek data dalam dua partisi.

- *Calinski–Harabasz index (CH)*, mengevaluasi validitas *cluster* berdasarkan jumlah rata-rata antara dan di dalam gugus kuadrat.
- *I index (I)*, mengukur pemisahan berdasarkan jarak maksimum antara pusat *cluster*, dan mengukur kekompakan berdasarkan jumlah jarak antara objek dan pusat *cluster*-nya.
- *Dunn's index (D)*, menggunakan jarak berpasangan minimum antara objek dalam kelompok yang berbeda sebagai pemisahan *intercluster* dan diameter maksimum di antara semua *cluster* sebagai kekompakan *intracluster*.
- *Silhouette index (S)*, menguji hasil kerja *clustering* yang didasarkan kepada perbandingan berpasangan jarak *intercluster* dan *intracluster*. Selain itu, jumlah kelompok terbaik ditentukan melalui cara memaksimalkan nilai indeks ini.
- *Davies–Bouldin index (DBI)*, kesamaan antara suatu kelompok K dan seluruh kelompok lainnya dihitung, dan nilai paling besar ditetapkan pada K sebagai kesamaan kelompoknya. Selanjutnya *DBI* bisa didapat melalui rata-rata seluruh kesamaan kelompok. Hasil pengelompokan semakin baik dan mendapatkan partisi paling baik jika indeks semakin kecil.
- *Xie-Beni index (XB)*, mendefinisikan pemisahan *intercluster* sebagai jarak kuadrat minimum antara pusat *cluster*, dan kekompakan *intracluster* sebagai jarak kuadrat rata-rata antara setiap objek data dan pusat *cluster*-nya. Jumlah *cluster* optimal tercapai ketika minimum *XB* ditemukan.
- *SD index (SD)*, ide dasarnya pada konsep hamburan rata-rata dan pemisahan total *cluster*. Ketentuan pertama mengevaluasi kekompakan berdasarkan varians objek *cluster*, dan ketentuan kedua mengevaluasi perbedaan pemisahan berdasarkan jarak antara pusat *cluster*. *SD index* adalah penjumlahan dari kedua ketentuan ini, dan jumlah *cluster* optimal dapat diperoleh dengan meminimalkan nilai *SD*.

- *S Dbw index (S Dbw)*, memperhitungkan kepadatan untuk mengukur pemisahan *intercluster*. Ide dasarnya adalah bahwa untuk setiap pasangan pusat *cluster*, setidaknya satu dari kepadatan mereka harus lebih besar dari kepadatan titik tengah mereka. Kekompakan *intracluster* sama dengan *SD*. Demikian pula, indeks adalah penjumlahan dari kedua istilah ini dan nilai minimum *S Dbw* menunjukkan nomor *cluster* optimal.
- *Clustering Validation index based on Nearest Neighbors (CV NN)*, mengevaluasi pemisahan *intercluster* berdasarkan objek yang membawa informasi geometris dari setiap *cluster*. CV NN menggunakan beberapa objek dinamis sebagai perwakilan untuk *cluster* yang berbeda dalam situasi yang berbeda ketika mengukur pemisahan *intercluster*. CV NN juga menggunakan jarak berpasangan rata-rata antara objek dalam *cluster* yang sama dengan pengukuran kekompakan *intracluster*. CV NN mengambil bentuk penjumlahan pemisahan *intercluster* dan kekompakan *intracluster* setelah normalisasi untuk keduanya.

Davies Bouldin Index merupakan salah satu metode yang digunakan untuk mengukur validitas *cluster* dalam metode pengelompokan, kohesi didefinisikan sebagai penjumlahan kedekatan data dengan titik pusat *cluster* dari *cluster* yang diikuti. *Davies Bouldin Index* termasuk dalam kategori *internal measures*. Hasil *clustering* terbaik ditentukan berdasarkan nilai *Davies Bouldin Index* yang mendekati 0 (Mughnyanti et al., 2020). Kelemahan umum untuk semua algoritma *cluster*, yaitu kinerjanya sangat bergantung pada pengguna yang mengatur berbagai parameter. Bahkan, pengaturan yang tepat biasanya hanya dapat ditentukan dengan metode *trial and error*. Secara substansial *Davies Bouldin Index* mengatasi kelemahan ini dengan hanya meminta pengguna untuk menentukan jarak dan ukuran dispersi yang akan digunakan (Davies & Bouldin, 1979). Untuk menghitung nilai *Davies Bouldin Index* perlu dihitung *Sum of Square Within Cluster (SSW)*, *Sum of Between Cluster (SSB)*, dan nilai rasio antar *cluster* (*R*) menggunakan persamaan berikut (Orisa, 2022):

$$SSW_1 = \frac{1}{m_i} \sum_{j=1}^{m_i} d(x_j, c_i) \quad (2.4)$$

m_i = jumlah data pada *cluster* ke-i

c_i = pusat *cluster* ke-i

$d(x_j, c_i)$ = jarak masing-masing data ke *centroid* (Euclidean)

$$SSB_{i,j} = D(C_i, C_j) \quad (2.5)$$

$$R_{ij} = \frac{SSW_1 + SSW_j}{SSB_{j,j}} \quad (2.6)$$

Nilai rasio digunakan untuk menghitung nilai *Davies Bouldin Index* dengan jumlah *cluster* yang telah ditentukan (K) menggunakan persamaan berikut:

$$DBI = \frac{1}{K} \sum_{K=1}^K \max_{i \neq j} (R_{(i,j)}) \quad (2.7)$$

Dummy Book Finish Wawasan Ilmu

BAB III

MODEL SISTEM INTELIJENSIA BISNIS

A. Model Sistem Intelijensia Bisnis di Pendidikan Tinggi

Universitas sebagai lembaga pendidikan tinggi berupaya untuk meningkatkan kompetensi dosen dengan cara mengevaluasi hasil penelitian mereka melalui publikasi ilmiah, hal ini dilakukan untuk meningkatkan standar pendidikan. Publikasi ilmiah adalah fase akhir dalam penyusunan artikel hasil penelitian, dimaksudkan untuk menyebarkan temuan kepada masyarakat luas. Dengan menerbitkan karya ilmiah, seorang akademisi juga ikut berpartisipasi dalam mengatasi masalah yang belum memiliki solusi. Selain itu, penerbitan karya ilmiah memiliki manfaat konkret bagi penulis, yang dapat menjadi pendorong untuk terus menulis dan mempublikasikan tulisan. Publikasi adalah elemen penting dalam meningkatkan karir dan seringkali diikuti dengan pengakuan atau imbalan baik secara langsung maupun tidak langsung.

Dalam menjalankan penelitian berkualitas, universitas menetapkan beberapa indikator untuk mengukur pencapaian standar hasil penelitian, termasuk: Adanya penelitian yang disebarkan melalui publikasi di jurnal nasional yang memiliki

akreditasi; Terdapat hasil penelitian yang diseminasi di jurnal internasional yang diakui reputasinya; Terdapat pengutipan (*citations*) terhadap publikasi para dosen yang relevan dengan program studi; Terdapat penelitian dosen yang sesuai dengan rencana pengembangan penelitian (*roadmap* penelitian).

Pengukuran kinerja publikasi ilmiah dosen di universitas ditentukan oleh pimpinan perguruan tinggi berdasarkan data jumlah karya ilmiah yang telah diterbitkan oleh dosen tersebut. Informasi mengenai publikasi ilmiah diperoleh dari lembaga pengindeks seperti *Science and Technology Index* (Sinta), Scopus, Web of Science, *Directory of Open Access Journal* (DOAJ), Proquest, ASEAN Citation Index (ACI), Google Scholar, dan *Indonesian Publication Index* (portalgaruda.org). Salah satu contoh situs web yang digunakan dalam pengukuran ini adalah Sinta, yang merupakan portal informasi penelitian yang digunakan untuk menilai kinerja peneliti, lembaga, dan jurnal di Indonesia. Sinta dikelola oleh Kementerian Pendidikan, Kebudayaan, Riset, dan Teknologi Republik Indonesia (Kemristekdikti RI). Dalam mengatur tata kelola perguruan tinggi, penting untuk memiliki sistem penyimpanan data yang dapat dikustomisasi sesuai dengan kebutuhan internal institusi, termasuk dalam proses pengukuran kinerja publikasi ilmiah dosen.

Informasi yang diperoleh dari situs web Sinta merupakan data yang berasal dari sumber eksternal, sehingga keterbatasan akses dan analisis data untuk keperluan internal institusi. Selain itu, pemanfaatan informasi mengenai publikasi ilmiah dosen belum dimaksimalkan dalam mengembangkan potensi dosen sesuai dengan karakteristik masing-masing. Untuk meningkatkan kualitas sumber daya dosen dalam bidang penelitian, akan dilakukan desain model *data warehouse* untuk menyimpan data publikasi ilmiah yang bersumber dari situs web Sinta dalam Sistem Intelijensia Bisnis. Selain itu, diterapkan teknik data mining dalam mengelompokkan peneliti guna mengukur kinerja publikasi ilmiah sebagai dasar untuk mengambil tindakan optimal dalam mengembangkan potensi dosen.

Dalam menyelesaikan masalah, dilakukan desain *data warehouse* dengan menggunakan model dimensional. Model multidimensi (dimensional) menggambarkan data dalam ruang n-dimensi yang

ditentukan oleh dimensi dan fakta. Dimensi adalah sudut pandang yang digunakan untuk menelaah data (Vaisman & Zimányi, 2014). Ada dua jenis model dimensional yang dipakai, pertama adalah model *star schema* yang terdiri dari kumpulan tabel dimensi dan tabel fakta. Kedua, terdapat model *snowflakes schema* yang terdiri dari tabel dimensi, sub dimensi, dan tabel fakta (Iqbal et al., 2020). Hasil dari penerapan query pada kedua teknik pemodelan dibandingkan dengan eksperimen pada data dan menghasilkan hasil yang serupa.

Pendekatan yang digunakan dalam *data mining* melibatkan penggunaan algoritma K-Means dan X-Means untuk mengelompokkan peneliti berdasarkan nilai *H-Index* mereka di platform Scopus, Google Scholar, dan Web of Science yang diakses melalui situs web Sinta. Metode pengelompokan K-Means adalah salah satu algoritma yang paling umum digunakan dalam menyelesaikan masalah pengelompokan, tetapi memiliki kesulitan dalam menentukan jumlah *cluster* secara apriori (Sinaga & Yang, 2020). X-Means *clustering* merupakan sebuah algoritma pengelompokan yang merupakan penyempurnaan dari K-means clustering, dengan tujuan untuk mengatasi masalah dalam menentukan jumlah *cluster* serta meningkatkan kecepatan proses pada dataset yang memiliki ukuran besar (Yusuf et al., 2022). Penentuan algoritma terunggul antara K-Means dan X-Means bergantung pada pengukuran *Davies Bouldin Index* (DBI).

Kubus *Online Analytical Processing* (OLAP) digunakan sebagai metode analisis data untuk memenuhi kebutuhan fungsional dalam memberikan informasi yang diperlukan. Sistem ini menghasilkan informasi yang disajikan dalam bentuk visualisasi data pada *dashboard* kecerdasan bisnis. Pengertian visualisasi data adalah menggunakan representasi grafis untuk memahami, mengeksplorasi, dan berkomunikasi dengan data (Sharda et al., 2018). Evaluasi sistem dilaksanakan untuk mengevaluasi sistem sebelum digunakan, dan terdapat dua kegiatan utama dalam evaluasi sistem kecerdasan bisnis, yaitu verifikasi dan validasi.

1. Analisis Kebutuhan Sistem BI Publikasi Ilmiah Dosen

Analisis kebutuhan bisnis pada sistem intelijensia bisnis merepresentasikan pertanyaan bisnis yang perlu dijawab oleh

pengambil keputusan bisnis (aktivitas kerja yang membutuhkan pengetahuan dan penilaian) untuk membuat keputusan yang efektif. Kerangka PIECES (*Performance, Information, Economy, Control, Efficiency, Services*) digunakan untuk mengkategorikan masalah, peluang, dan arahan dalam menganalisis serta merancang sistem, yang menjadi pedoman penting dalam pengembangan sistem informasi (Setiyani et al., 2020). Melalui analisis PIECES dapat diketahui perbedaan antara sebelum penerapan dan sesudah penerapan Sistem Intelijensia Bisnis (Hidayat & Fitriana, 2022). Tabel 4 menunjukkan contoh hasil analisis kebutuhan pada Sistem Intelijensia Bisnis berdasarkan pendekatan PIECES.

Tabel 8. Matriks Hasil Analisis PIECES

PIECES	Sebelum Ada Sistem BI	Setelah Ada Sistem BI
Performance: Sistem mampu menyelesaikan tugas dengan cepat, memungkinkan pencapaian target secara efisien.	Proses pencarian data publikasi ilmiah membutuhkan waktu yang relatif lama.	Melalui sistem <i>business intelligence</i> , pencarian data publikasi ilmiah dapat dilakukan dengan efisien.
Information: Sistem dapat menghasilkan informasi yang akurat sesuai dengan kebutuhan yang ada.	Terdapat keterbatasan dalam akses dan analisis data untuk keperluan internal institusi.	Informasi dan analisis data untuk keperluan internal institusi dapat diakses sesuai dengan kebutuhan.
Economy: Penerapan sistem memberikan keuntungan dalam penyediaan informasi dengan biaya yang ekonomis.	Pengolahan data publikasi ilmiah dilakukan secara manual oleh operator, sehingga memerlukan alokasi biaya untuk sumber daya manusia.	Pengurangan biaya sumber daya manusia sebagai operator dapat dicapai karena sistem secara otomatis mengakses dan memproses data.

Control: Sistem memiliki keunggulan dalam memastikan akses mudah, integritas, dan keamanan data.	Terdapat kurangnya konsistensi dalam data karena belum terintegrasi dengan baik.	Penggunaan model dimensional dalam sistem business intelligence memastikan integritas referensial sebagai upaya untuk menjaga konsistensi data.
Efficiency: Penerapan sistem berdampak pada penggunaan sumber daya yang lebih efisien.	Proses penyusunan laporan membutuhkan waktu dan tenaga yang tidak optimal.	Proses penyusunan laporan dapat dilakukan dengan cepat menggunakan visualisasi data.
Services: Penerapan sistem menghasilkan peningkatan layanan yang lebih baik bagi organisasi.	Keterbatasan akses dan analisis data, serta pengolahan data publikasi ilmiah secara manual, dapat mengakibatkan pengambilan keputusan yang terhambat.	Pengambilan keputusan dapat dilakukan dengan efisien berdasarkan informasi yang tersedia dalam sistem <i>business intelligence</i> .

Proses pengembangan sistem kecerdasan bisnis dimulai dengan mengidentifikasi kebutuhan fungsional, yang melibatkan penentuan fitur-fitur yang diperlukan oleh pengguna sebagai fondasi utama (Setiyani et al., 2020). Kebutuhan fungsional dari sistem kecerdasan bisnis untuk menilai kinerja publikasi ilmiah berdasarkan hasil analisis adalah sebagai berikut:

- Memberikan informasi mengenai indeks peneliti berdasarkan indeks, dosen, dan program studi.
- Memberikan informasi mengenai nilai peneliti berdasarkan dosen, program studi, dan *cluster* peneliti.
- Memberikan informasi mengenai artikel publikasi berdasarkan

indeks, tahun penerbitan, dosen, dan program studi.

- Memberikan informasi mengenai subjek penelitian berdasarkan dosen, program studi, dan rencana pengembangan penelitian.

Pentingnya sistem intelijensia bisnis adalah untuk mengubah data-data menjadi informasi yang bermakna dan bermanfaat bagi perusahaan. Secara umum, dalam pengembangan sistem intelijen bisnis, terdapat serangkaian tahapan yang dilalui untuk mencapai efektivitas yang diinginkan, yaitu:

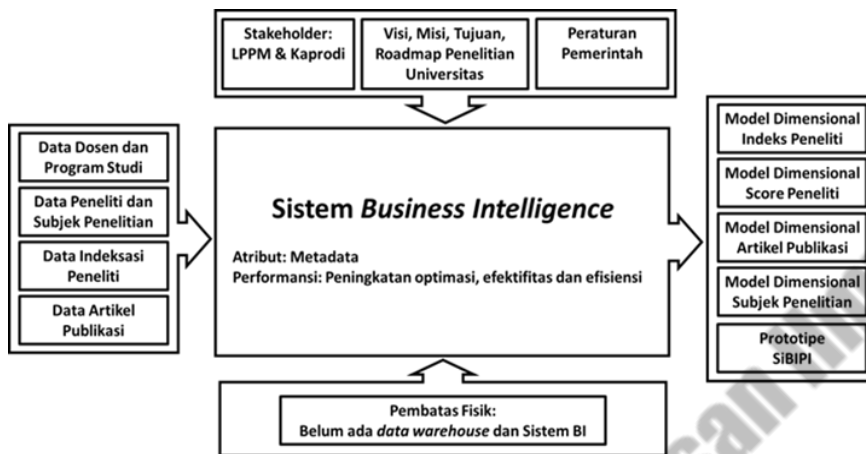
- **Analisis Kebutuhan Data dan Informasi:** Melibatkan pengidentifikasian informasi yang diperlukan untuk merancang sistem intelijensia bisnis. Proses ini kemudian dilanjutkan dengan pengumpulan data sesuai dengan hasil analisis kebutuhan data tersebut.
- **Perancangan Data Warehouse:** Tahap perencanaan *data warehouse* difokuskan pada menciptakan struktur dan rencana dasar untuk data warehouse tersebut. Dalam tahap ini, akan ditetapkan komponen yang diperlukan agar setiap elemen dalam *data warehouse* dapat saling terhubung dan dimanfaatkan secara efektif saat digunakan.
- **Penerapan Data Mining:** Teknik *clustering* adalah bentuk pembelajaran mesin di mana data tidak disertai dengan label atau informasi yang ditentukan sebelumnya. Teknik *clustering* K-Means dan X-Means digunakan untuk memisahkan sekelompok data menjadi beberapa bagian atau kategori yang sesuai.
- **Akuisisi Data:** Proses ekstraksi informasi bisnis yang relevan dari berbagai sistem sumber operasional serta mengubah data menjadi format homogen dan memuat data ke *Data Warehouse*.
- **Penerapan OLAP:** Pemanfaatan OLAP dioptimalkan dengan menginterpretasi data guna memahami peristiwa beserta akar penyebabnya, mengenali potensi peristiwa di masa depan, serta mengidentifikasi situasi terkini untuk membuat keputusan yang menghasilkan kinerja optimal.

- **Pengembangan Sistem BI:** Menggabungkan elemen inti dari sistem, yaitu *data warehouse* sebagai sumber informasi yang akan diproses oleh *business analytics* untuk menghasilkan data mengenai aktivitas yang mencerminkan kinerja bisnis untuk pemantauan dan analisis yang disajikan melalui antarmuka pengguna berupa visualisasi data dalam bentuk *dashboard*, sesuai dengan struktur sistem kecerdasan bisnis yang telah direncanakan.
- **Evaluasi Sistem:** Terdapat dua kegiatan utama yang dilakukan, yaitu verifikasi dan validasi. Verifikasi merupakan proses untuk membuktikan bahwa sistem beroperasi tanpa kesalahan dan sesuai dengan spesifikasi yang telah ditetapkan. Sementara itu, uji validasi dilakukan untuk menunjukkan bahwa sistem dapat memberikan informasi sesuai dengan harapan dan kebutuhan pengguna.

Diagram alir metode yang menggambarkan urutan langkah-langkah dalam perancangan model sistem intelijensia bisnis dapat dilihat pada Gambar 14.

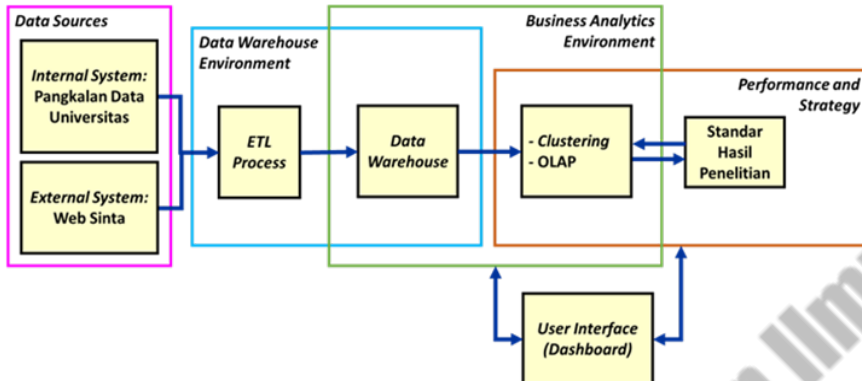
2. Model Sistem BI Publikasi Ilmiah Dosen

Sistem kecerdasan bisnis yang telah dirancang mencakup elemen inti dari sistem yang terhubung, yang mencakup *data warehouse* sebagai sumber data yang dianalisis oleh *business analytics*. Hal ini bertujuan untuk menghasilkan informasi mengenai aktivitas publikasi ilmiah yang sesuai dengan manajemen kinerja bisnis untuk memantau dan mengevaluasi kinerja melalui antarmuka pengguna berupa visualisasi data dalam bentuk *dashboard*.



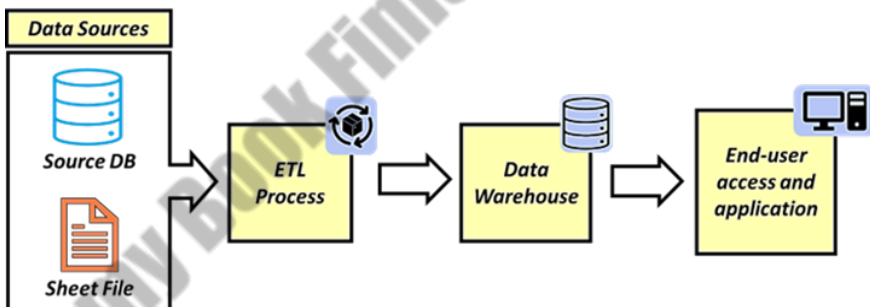
Gambar 14. Model Sistem BI Publikasi Ilmiah Dosen

Gambar 13 memperlihatkan bahwa input untuk Sistem *Business Intelligence* Publikasi Ilmiah (SiBIPi) terdiri dari informasi mengenai dosen dan program studi, data peneliti dan subjek penelitian, data indeksasi peneliti, dan data artikel publikasi. Namun, terdapat kendala fisik karena belum tersedianya sistem kecerdasan bisnis dan *data warehouse*. *Stakeholder* yang menggunakan sistem ini adalah staf dari Lembaga Penelitian dan Pengabdian Masyarakat (LPPM) serta ketua program studi. Sistem ini dipengaruhi oleh visi, misi, tujuan, dan rencana pengembangan penelitian universitas serta peraturan yang ditetapkan oleh pemerintah. Karakteristik dari sistem kecerdasan bisnis ini mencakup atribut dari model data dimensional. Tujuan utama performansi dari sistem ini adalah meningkatkan efisiensi, efektivitas, dan optimalisasi. Hasil dari sistem ini adalah informasi mengenai indeks peneliti, nilai peneliti, artikel publikasi, dan subjek penelitian. Hasil akhir dari proyek ini adalah prototipe dari Sistem *Business Intelligence* Publikasi Ilmiah (SiBIPi). Model arsitektur dari sistem kecerdasan bisnis yang telah dirancang diilustrasikan dalam Gambar 14.



Gambar 15. Model Arsitektur Sistem BI Publikasi Ilmiah Dosen

Perancangan *data warehouse* melibatkan beberapa tahapan, termasuk perancangan model data, transformasi data dari sumber ke tujuan (ETL), dan pemilihan model terbaik. Sistem arsitektur *data warehouse* yang diterapkan dalam desain ini adalah arsitektur terpusat. Ilustrasi dari model arsitektur *data warehouse* untuk publikasi ilmiah dapat ditemukan pada Gambar 15.



Gambar 16. Model Arsitektur Data Warehouse Publikasi Ilmiah Dosen

Data berasal dari sumber-sumber seperti basis data dan file lembar kerja elektronik. Proses ETL digunakan untuk mengekstraksi dan mentransformasi data dari sumber tersebut, kemudian data dimuat ke dalam sistem *data warehouse*. Data yang telah disimpan di *data warehouse* dapat diakses oleh pengguna melalui aplikasi khusus.

Desain model data melibatkan penerapan dua model dimensional, yakni model star schema dan model snowflakes schema. Terdapat empat tahapan dalam proses perancangan model dimensional (Sugiarto et al., 2021), yang mencakup:

1. Memilih proses bisnis.
2. Menetapkan tingkat detail (*grain*).
3. Mengidentifikasi dimensi.
4. Mengidentifikasi fakta.

Hasil dari menetapkan tingkat detail dan mengidentifikasi dimensi untuk proses bisnis pengukuran kinerja publikasi ilmiah adalah hubungan antara tingkat detail dan dimensi. Tabel fakta dibuat berdasarkan hubungan tingkat detail dengan dimensi. Pada Tabel 4, dapat dilihat hubungan antara tingkat detail dan dimensi.

Tabel 9. Relasi Tingkat Detail (Grain) dan Dimensi

Dimensi Grain	Dosen	Peng indeks	Prodi	Tahun	Road map	Cen troid
Indeks	✓	✓	✓			
Score	✓		✓			✓
Artikel	✓	✓	✓	✓		
Subjek	✓		✓		✓	

Pada desain model *data warehouse* untuk publikasi ilmiah, *grain* mencakup:

- Indeks, menyimpan data mengenai indeks peneliti berdasarkan indeks, dosen, dan program studi.
- *Score*, menyimpan informasi mengenai nilai peneliti berdasarkan dosen, program studi, dan *cluster* peneliti.
- Artikel, menyimpan data mengenai artikel publikasi berdasarkan indeks, tahun penerbitan, dosen, dan program studi.

- Subjek, menyimpan informasi mengenai subjek penelitian berdasarkan dosen, program studi, dan rencana pengembangan penelitian.

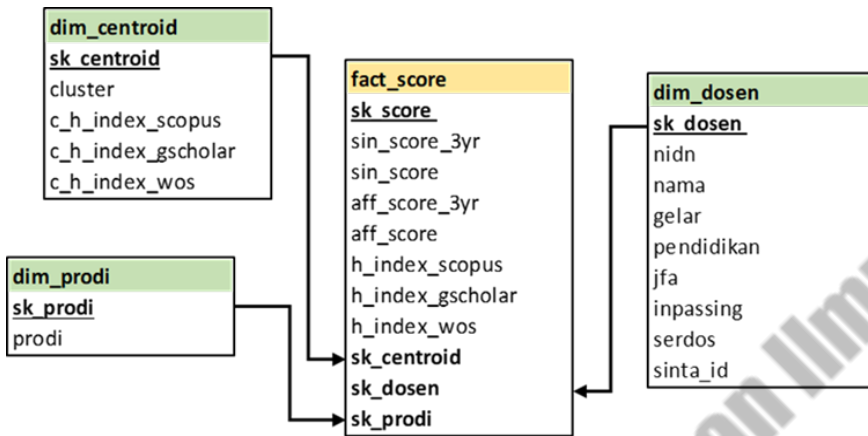
Elemen-elemen dalam desain model *data warehouse* untuk publikasi ilmiah, mencakup:

- Dosen, berisi informasi mengenai profil para dosen.
- Pengindeks, memuat data mengenai nama-nama pengindeks jurnal ilmiah.
- Prodi, berisi informasi mengenai nama-nama program studi.
- Tahun, memuat data tentang rentang tahun.
- *Roadmap*, menyimpan informasi mengenai subjek penelitian dalam roadmap penelitian.
- *Centroid*, berisi informasi mengenai nilai centroid dari masing-masing kelompok.

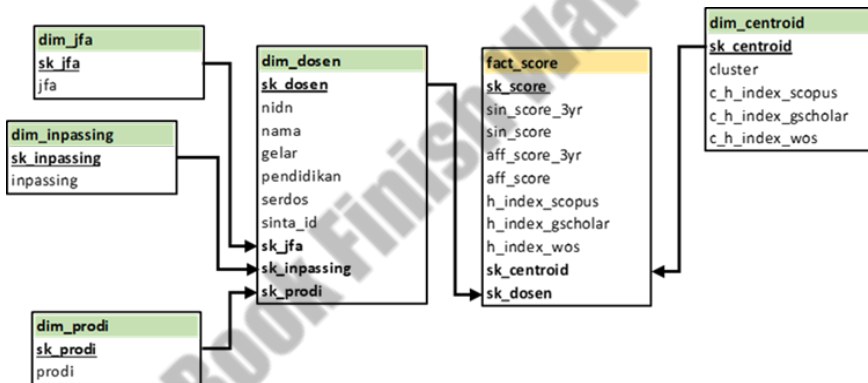
Dalam model *snowflakes schema*, terdapat tabel subdimensi yang terhubung dengan tabel dimensi. Subdimensi yang terbentuk dalam model skema salju mencakup:

- JFA, berisi informasi tentang jenjang jabatan fungsional akademik dosen.
- Inpassing, memuat data tentang jenjang inpassing dosen.
- Prodi, berisi informasi mengenai nama-nama program studi.

Salah satu model dimensional yang telah dibuat adalah model dimensional *score* peneliti yang menghasilkan informasi terkait *score* peneliti berdasarkan dosen, program studi, dan pengelompokan peneliti. Informasi terkait *score* peneliti dapat dijelaskan melalui hubungan antara fakta *score*, dimensi dosen, dimensi prodi, dan dimensi *centroid*. Model dimensional *score* peneliti tersebut tergambar pada Gambar 16 dan Gambar 17 menggunakan model *star schema* dan *snowflakes schema*.



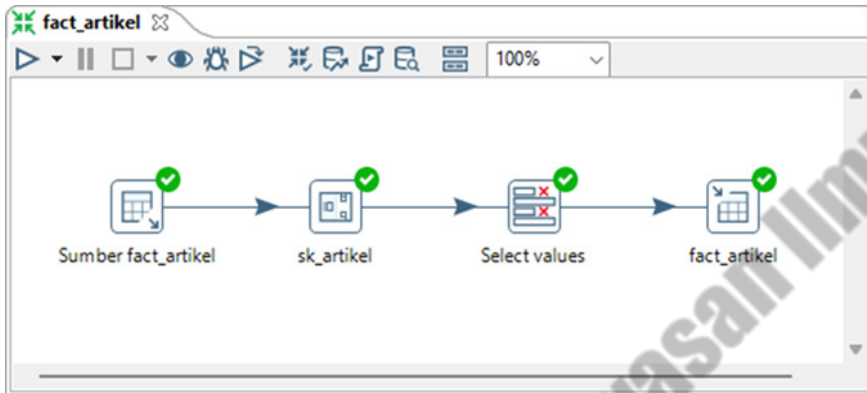
Gambar 17. Model Dimensional Score Peneliti Menggunakan Star Schema



Gambar 18. Model Dimensional Score Peneliti Menggunakan Snowflakes Schema

Proses *Extract, Transform, Load* (ETL) diperlukan ketika data akan dimuat ke dalam sistem *data warehouse* oleh *load manager*. Proses ini memastikan bahwa data yang dimasukkan ke dalam *data warehouse* telah memenuhi standar yang telah ditetapkan. Proses ETL untuk tabel dimensi dan tabel fakta dalam model *Star Schema* dan model *Snowflakes Schema* dijalankan menggunakan perangkat lunak Pentaho Data Integration (Pentaho). Sebagai contoh, proses ETL untuk tabel fakta artikel (*fact_artikel*) pada Pentaho menggunakan

serangkaian komponen *step* yang terdiri dari *table input*, *add sequence*, *select values*, *table output*. Hubungan antar komponen *step* tersebut dapat dilihat dalam Gambar 18.



Gambar 19. Komponen Step pada ETL fact_artikel

- *Table input step* digunakan untuk mengambil data yang diperlukan untuk proses transformasi. Data yang diperlukan untuk membentuk fact_artikel diambil dari beberapa tabel dalam database sumber menggunakan perintah SQL.
- *Add sequence step* digunakan untuk menambahkan *surrogate key* (SK) yang berperan sebagai kunci primer untuk mengidentifikasi setiap baris data. Ini juga memastikan konsistensi analisis jika terjadi perbedaan kunci pada entitas yang sama tetapi berasal dari sumber data yang berbeda.
- *Select values step* digunakan untuk memilih *field* yang diperlukan, menentukan properti *field*, dan mengatur urutan *field*.
- *Table output step* digunakan untuk memuat hasil transformasi ke dalam *data warehouse*.

Pada model *star schema* dan *snowflakes schema*, proses transformasi menggunakan perangkat lunak Pentaho diatur dalam suatu alur kontrol yang menjalankan seluruh proses transformasi secara berurutan dan dapat dijadwalkan.

Data *data warehouse* data disimpan dalam sebuah basis data. MySQL digunakan sebagai sistem manajemen basis data di *data*

warehouse, dan dikelola melalui antarmuka grafis SQLyog. Data dalam model *star schema* dan model *snowflakes schema* disimpan dalam basis data yang terpisah. Informasi mengenai basis data model *star schema* terdapat pada Tabel 5, sedangkan informasi mengenai basis data model *snowflakes schema* terdapat pada Tabel 6.

Tabel 10. Database Model Star Schema *Database: publikasi_star*

Name	Rows	Data Size (KB)
dim_centroid	1	16
dim_dosen	34	16
dim_pengindeks	3	16
dim_prodi	4	16
dim_roadmap	7	16
dim_tahun	100	16
fact_artikel	760	192
fact_indeks	102	16
fact_score	34	16
fact_subjek	56	16
Database Size (KB)		336

Tabel 11. Database Model Snowflakes Schema *Database: publikasi_snow*

Name	Rows	Data Size (KB)
dim_centroid	1	16
dim_dosen	34	16
dim_inpassing	5	16
dim_jfa	4	16
dim_pengindeks	3	16
dim_prodi	4	16
dim_roadmap	7	16
dim_tahun	100	16

fact_artikel	760	192
fact_indeks	102	16
fact_score	34	16
fact_subjek	56	16
Total Size (KB)		368

Eksperimen untuk menguji kedua model dilakukan pada data dan hasil yang sama dari menerapkan *query* pada kedua model dimensional. Tujuannya adalah untuk menentukan model mana, antara *star schema* dan *snowflakes schema*, yang memberikan kinerja yang lebih baik dalam *data warehouse* untuk kasus ini. Hasil pengujian *query* model *star schema* dapat dilihat pada Tabel 9 dan hasil pengujian *query* model *snowflakes schema* dapat dilihat pada Tabel 10.

Tabel 12. Waktu Proses Query Model Star Schema

No	Eksperimen Query	Waktu Proses (sec)
1	Informasi tentang indeks.	0,00072
2	Informasi tentang <i>score</i> .	0,00059
3	Informasi tentang artikel publikasi.	0,00363
4	Informasi tentang subjek penelitian.	0,00060
	Total Waktu Proses	0,00554

Tabel 13. Waktu Proses Query Model Snowflakes Schema

No	Eksperimen Query	Waktu Proses (sec)
1	Informasi tentang indeks.	0,00089
2	Informasi tentang <i>score</i> .	0,00075
3	Informasi tentang artikel publikasi.	0,00371
4	Informasi tentang subjek penelitian.	0,00076

	Total Waktu Proses	0,00611
--	--------------------	---------

Perbandingan hasil pengujian dapat dilihat pada Tabel 10.

Tabel 14. Perbandingan Model Dimensional

Model Dimensional	Total Size (KB)	Waktu Proses (sec)
<i>Star schema</i>	336	0,00554
<i>Snowflakes schema</i>	368	0,00611

Berdasarkan hasil pengujian *query* pada model *star schema* dan *snowflakes schema* menunjukkan bahwa model *star schema* adalah model terbaik yang dipilih, karena memiliki ukuran data lebih kecil sebesar 336 KB dan waktu pemrosesan *query* lebih singkat sebesar 0,00554 *second*.

B. Studi Kasus di Industri Provider

Indonesia sebagai salah satu negara berkembang dengan potensi pasar yang sangat besar, mengawali era penggunaan *internet* di awal tahun 1990-an. Pada saat itu, jaringan *internet* lebih populer dengan sebutan paguyuban *network* yang berdasarkan kerjasama, kekeluargaan, dan gotong royong. Suatu kondisi yang cukup berbeda dibandingkan saat ini dimana penggunaan *internet* lebih mengutamakan aspek bisnis dan mencari keuntungan.

Saat ini, penggunaan *internet* di Indonesia tidak lagi sama seperti awal kehadirannya. Hampir seluruh aspek kehidupan melibatkan peran *internet* dalam kehidupan masyarakat. Metode pembayaran seperti tagihan, parkir, pembayaran belanja, hingga langganan berbagai layanan bisa dilakukan dengan mudah dengan keberadaan *internet* di *smartphone*. Transfer dan transaksi antar bank yang dahulu harus ke bank ataupun ATM, kini bisa dilakukan dimana saja selama ada akses *internet*. Hal tersebut semakin dipermudah dengan kehadiran teknologi 4 *Generations* (4G) serta pengembangan teknologi 5 *Generations* (5G).

Dalam sebuah penelitian yang dilakukan sebelumnya diperoleh data kebiasaan penggunaan *internet* di Indonesia, dimana 36,8% pengguna menghabiskan 7-9 jam dan 34,4% pengguna menghabiskan lebih dari 9 jam dalam menggunakan *internet*. Dalam memanfaatkan akses *internet*, 58% responden memanfaatkan penggunaan *smartphone* dan 41,1% menggunakan lebih dari satu *gadget* (Prasetyo et al. 2021).

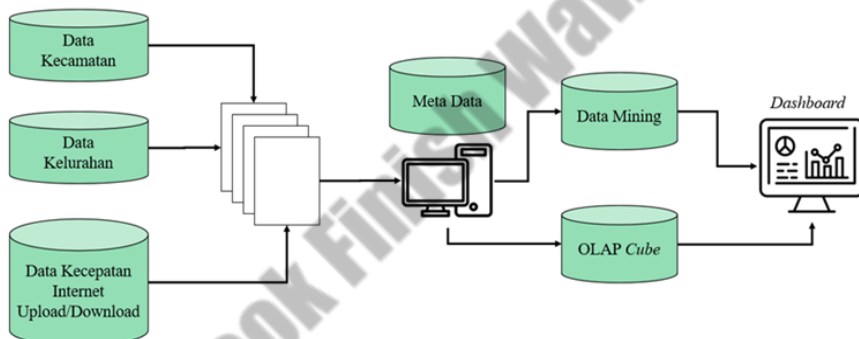
Akses kemudahan *internet* juga menjadi faktor berkembangnya teknologi *blockchain*. *Blockchain* adalah sebuah teknologi untuk distribusi terdesentralisasi, dimana secara teknis dapat mengatasi masalah keamanan yang muncul pada keyakinan terhadap model sentralisasi. Lebih jauh, para peneliti mengkombinasikan *blockchain* dan kontrol akses pada proteksi data *internet* (2). Hal ini juga didukung oleh penelitian lain dimana disebutkan bahwa *blockchain* mampu menjamin keamanan pertukaran data dan memungkinkan komunitas *digital* untuk terhubung dan membentuk suatu jaringan (Thomason 2022).

Perkembangan *Internet of Things* (IoT) memberikan suatu kesempatan bagi bisnis-bisnis baru di samping mendisrupsi beberapa bisnis yang sudah eksis (Leiting, De Cuyper, and Kauffmann 2022). Dimana perkembangan IoT tidak lepas dari semakin luasnya jangkauan akses *internet* di dunia. IoT sendiri merupakan sebuah infrastruktur dari entitas, manusia, sistem, dan sumber daya informasi yang saling berhubungan bersama-sama dengan layanan yang memproses dan bereaksi terhadap informasi dari dunia fisik dan dunia maya (Aloraini et al. 2022). Hal ini menjadi kesempatan baru bagi perusahaan operator seluler untuk meningkatkan profitabilitas dari semakin beragamnya penggunaan data *internet* di masyarakat.

Penerapan sistem intelijensia bisnis adalah untuk memperoleh informasi yang tersembunyi dari data-data yang dimiliki dan memiliki makna bagi suatu perusahaan/organisasi. Pada penelitian in, penerapan sistem intelijensia bisnis dilakukan dengan beberapa tahapan yang dilakukan, seperti yang ditunjukkan pada gambar berikut ini:

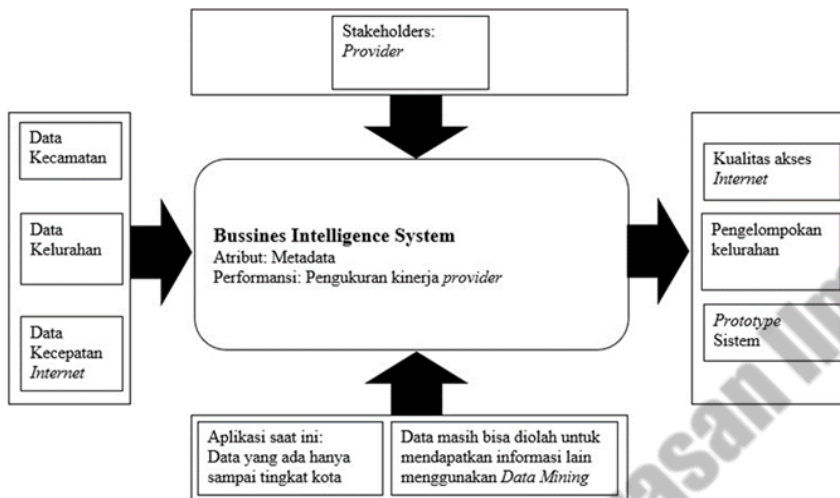
1. Model Sistem Intelijensia Bisnis Kecepatan Provider

Penelitian yang dilakukan bertujuan untuk melakukan perancangan *BI* untuk informasi kecepatan *internet* dari setiap operator GSM di setiap kelurahan kota Bekasi. Data yang didapatkan selanjutnya disatukan ke dalam 1 data *warehouse* untuk menghasilkan suatu tabel fakta yang akan diolah menggunakan *OLAP*. Data yang diperoleh selanjutnya digunakan untuk proses *clustering* kelurahan di kota Bekasi berdasarkan kecepatan akses *internet* baik *upload* maupun *download* menggunakan metode *K-Means* dan *K-Medoids* untuk selanjutnya dibandingkan akurasi dari hasil yang didapatkan. Algoritma dengan nilai *DBI* terkecil dianggap sebagai algoritma terbaik dengan nilai *k* yang telah ditentukan. Gambar 20 menunjukkan model arsitektur dari sistem *BI* yang akan dibuat pada penelitian ini.



Gambar 20. Model arsitektur BI

Penelitian ini menggunakan sistem *Business Intelligence* untuk mendapatkan informasi mengenai kecepatan *internet* di setiap kelurahan kota Bekasi. dalam proses pembuatan sistem *BI* yang akan dibuat, langkah awal yang perlu dilakukan adalah menentukan informasi *input* yang dibutuhkan dan informasi *output* yang akan dibuat berdasarkan beberapa kondisi saat ini serta penentuan performansi yang diharapkan (Fitriana, Saragih, and Hasyati 2018). Penggunaan sistem *BI* ini nantinya akan didukung oleh tampilan *dashboard* yang menampilkan informasi-informasi yang dibutuhkan. Gambar 21 menunjukkan *input* hingga *output* pada sistem *BI* pada peneltian ini.



Gambar 21. Sistem BI

Gambar 21 merupakan bentuk sistem dari rancangan sistem yang akan dibuat pada penelitian ini. Data kelurahan dan data kecamatan menjadi data sekunder yang akan digunakan dalam penelitian ini. Data kecepatan *internet upload* dan *download* merupakan data primer yang digunakan untuk membentuk sistem BI yang akan dibuat. Data tersebut kemudian dikumpulkan ke dalam *data warehouse* melalui proses ETL. Langkah selanjutnya adalah mengolah informasi tersebut untuk menghasilkan informasi baru menggunakan OLAP dan *data mining* yang selanjutnya akan dimunculkan pada sebuah *dashboard*.

2. Proses ETL

Penelitian ini melalui tahap ETL untuk melakukan proses ekstraksi, transformasi, dan memasukkan data ke dalam *database* sehingga berisi tabel dimensi dan tabel fakta yang akan digunakan. Proses ETL pada penelitian ini dilakukan dengan bantuan aplikasi *Pentaho*.

3. Proses ETL untuk Tabel Dimensi

Penelitian ini membutuhkan beberapa tabel dimensi untuk membentuk *data warehouse* yang dibutuhkan, antara lain tabel dimensi kecamatan (*dim_kecamatan*), tabel dimensi kelurahan

(dim_kelurahan), tabel dimensi *upload* (dim_upload), dan tabel dimensi *download* (dim_download).

Tabel dim_kelurahan dihasilkan dengan melakukan proses *ETL* pada data kelurahan yang didapatkan melalui proses *data scrapping*. Data kelurahan yang didapatkan terdiri dari beberapa atribut yaitu kode kemendagri, kecamatan, jumlah kelurahan, dan daftar kelurahan. Gambar 22 menunjukkan proses *ETL* yang dilakukan menggunakan *Pentaho*.



Gambar 22. Proses ETL pada tabel dim_kecamatan

Proses *ETL* yang ditunjukkan pada gambar 22 dimulai dari melakukan *input* data menggunakan operator *Microsoft Excel Input*. Data tersebut selanjutnya ditambahkan penomoran *increment* untuk mengurutkan setiap data yang ada menggunakan *Add Sequence*. Proses penambahan *increment* ini digunakan sebagai *Primary Key* untuk tabel dim_kecamatan. Proses berikutnya yaitu melakukan pemilihan dan penamaan atribut yang akan digunakan menggunakan operator *Select Values*. Atribut Kode Kemendagri selanjutnya diubah nama menjadi Kode Kecamatan. Tabel dim_kecamatan yang dibuat terdiri dari 3 atribut yaitu sk_kecamatan, Kode Kemendagri yang selanjutnya diubah menjadi Kode Kecamatan, dan Kecamatan. Atribut lainnya pada tabel data kecamatan tidak digunakan yaitu Jumlah Kelurahan dan Daftar Kelurahan. Tabel 11 merupakan tabel dim_kecamatan yang telah di *export* ke dalam *database internet* yang dibuat sebelumnya.

Tabel 15. Tabel dim_kecamatan

sk_kecamatan	Kode_Kecamatan	Kecamatan
1	32.75.07	Bantar Gebang
2	32.75.02	Bekasi Barat
3	32.75.04	Bekasi Selatan

4	32.75.01	Bekasi Timur
5	32.75.03	Bekasi Utara
6	32.75.09	Jatiasih
7	32.75.10	Jatisampurna
8	32.75.06	Medan Satria
9	32.75.11	Mustika Jaya
10	32.75.08	Pondok Gede
11	32.75.12	Pondok Melati
12	32.75.05	Rawalumbu

Tabel selanjutnya yang akan dibuat adalah tabel `dim_kelurahan` yang merupakan tabel dimensi untuk informasi tentang kelurahan yang ada di kota Bekasi. proses *ETL* untuk menghasilkan tabel dimensi ini ditunjukkan pada gambar 23.



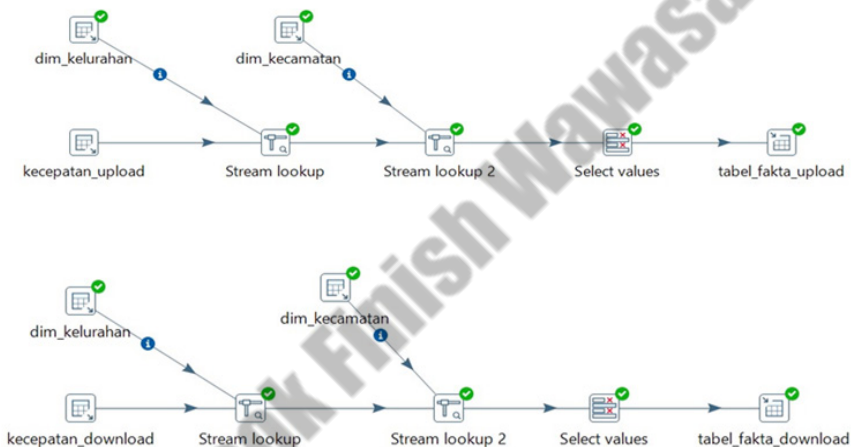
Gambar 23. Proses ETL pada tabel `dim_kelurahan`

Tabel `dim_kelurahan` didapatkan dari data kelurahan dan data kecamatan yang digabungkan dengan operator *Stream Lookup*. Kedua data tersebut digabungkan dengan tujuan mendapatkan data kode kecamatan untuk masing masing data kelurahan. Operator *Add Sequence* digunakan untuk mendapatkan atribut *Primary Key* yang akan digunakan untuk membentuk *star schema* yang akan dibuat. Tabel `dim_kelurahan` yang telah di *export* ke dalam *database internet* dapat dilihat pada bagian lampiran.

4. Proses Pembuatan Tabel Fakta

Penelitian ini menggunakan dua buah tabel fakta berdasarkan data primer yang didapatkan, yaitu `tabel_fakta_upload` dan `tabel_`

fakta_download. Gambar 24 memperlihatkan proses pembuatan tabel fakta baik untuk *data warehouse* kecepatan *upload* maupun kecepatan *download*. Dari gambar tersebut terlihat bahwa kedua tabel fakta menggunakan tabel dimensi *dim_kelurahan* dan *dim_kecamatan* untuk membentuk kedua tabel fakta tersebut. Perbedaan antara kedua tabel fakta tersebut terletak pada data kecepatan internet yang digunakan tergantung tabel fakta yang akan dibuat. Hasil output pada kedua tabel fakta tersebut adalah atribut *Kode_Kelurahan*, *Kode_Kecamatan*, *Kode_Pengambilan_Data*, *Provider*, dan *Kecepatan_upload* untuk *tabel_fakta_upload* serta *Kecepatan_download* untuk *tabel_fakta_download*.



Gambar 24. Proses pembentukan tabel fakta

5. Perancangan Data Warehouse

Penelitian ini menggunakan MySQL dalam merancang *database* sebagai tempat penyimpanan tabel dimensi dan tabel fakta yang akan digunakan. Tabel dimensi dan tabel fakta merupakan tabel yang akan membentuk skema *data warehouse* yang akan dibuat. Untuk membuat *database*, dibutuhkan beberapa aplikasi pendukung yaitu XAMPP dan MySQL.

Penelitian ini akan menggunakan *star schema data warehouse* untuk proses pembuatan *data warehouse*. Pembuatan *data warehouse*

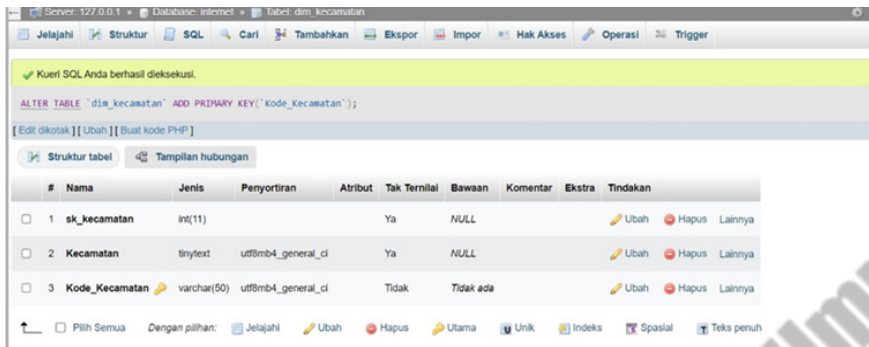
akan dilakukan setelah proses *ETL* selesai. Proses diawali dengan penentuan *Primary Key (PK)* pada tabel dimensi dan *Foreign Key (FK)* pada tabel fakta. Penelitian ini menggunakan 2 *data warehouse* yaitu *data warehouse* untuk data kecepatan *upload* dan *data warehouse* untuk kecepatan *download*. Gambar 24 menunjukkan skema *data warehouse* yang akan dibuat pada penelitian ini.

Langkah awal dalam penyusunan *data warehouse* yang dibutuhkan yaitu melakukan deklarasi *grain* yang akan menjadi dasar dalam menentukan tabel dimensi dan tabel fakta yang akan dibuat. Tabel 9 menunjukkan pemilihan *grain* terkait dengan penentuan informasi yang akan dimunculkan pada tabel fakta. Pada tabel tersebut ditentukan tabel fakta yang akan dibuat adalah informasi tentang kecepatan internet *upload* dan *download* dari setiap *provider*. *Data warehouse* yang akan dibuat pada penelitian ini menggunakan *star schema* yang terdiri dari tabel dimensi *dim_kecamatan* dan *dim_kelurahan*, serta tabel fakta untuk kecepatan *upload* dan *download*.

Tabel 16. Tabel grain dalam pembentukan data warehouse

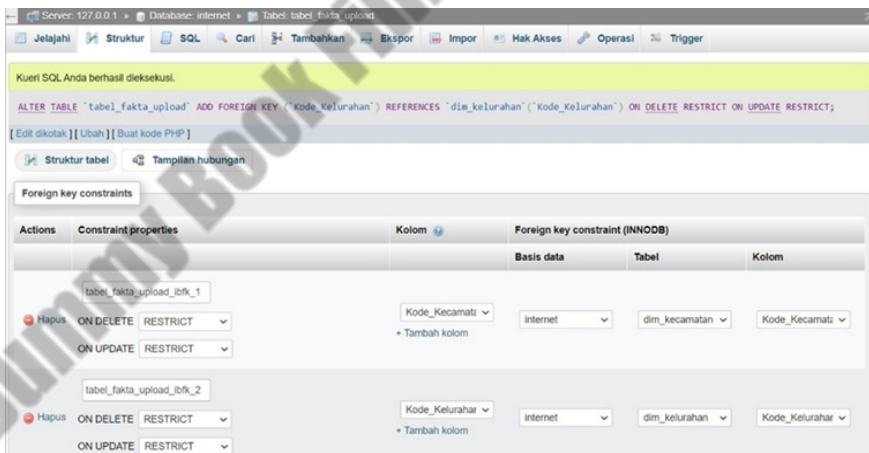
Dimensi	Grain	
	Kecepatan Upload	Kecepatan Download
Kecamatan	V	V
Kelurahan	V	V

Tabel dimensi dan tabel fakta yang akan digunakan selanjutnya digabungkan menjadi suatu skema yang disebut *star schema*. Pembentukan diawali dengan penentuan atribut yang menjadi *primary key* pada tabel dimensi dan *foreign key* pada tabel fakta. Gambar 25 menunjukkan cara penentuan *primary key* pada tabel dimensi yang akan digunakan. Pada tabel *dim_kecamatan*, atribut *Kode_Kecamatan* yang akan dijadikan *primary key* dipilih dan dijadikan utama seperti pada gambar. Hal serupa juga dilakukan untuk tabel *dim_kelurahan*



Gambar 25. Proses penentuan primary key pada tabel dimensi

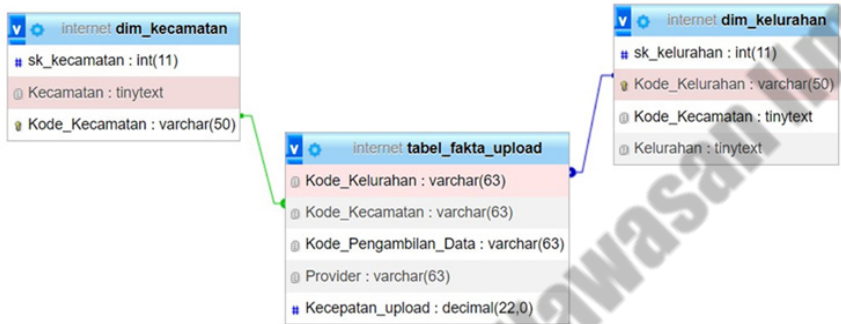
Penentuan *foreign key* pada tabel fakta dilakukan untuk menghubungkan atribut yang sama pada tabel dimensi dan tabel fakta. Proses diawali dengan memilih tabel fakta upload dan pilih struktur. Langkah selanjutnya adalah menuliskan *query* yang berfungsi untuk menentukan atribut yang akan dijadikan *foreign key*, yaitu tabel Kode_Kecamatan dan Kode_Kelurahan. Gambar 26 berikut ini menunjukkan *query* yang dituliskan dan hasil yang didapatkan untuk tabel_fakta_upload.



Gambar 26. Proses penentuan foreign key

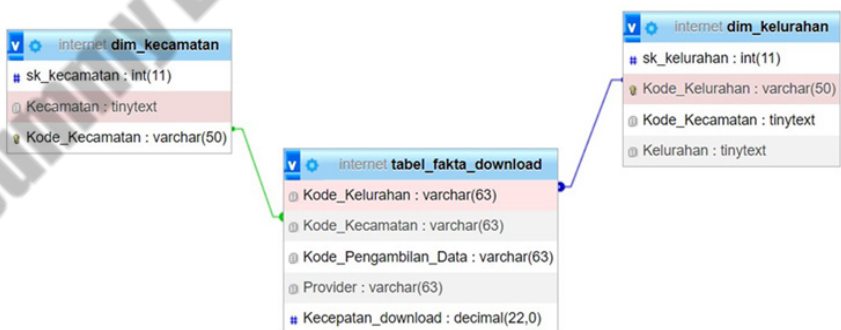
Penentuan *primary key* dan *foreign key* pada proses sebelumnya digunakan untuk menghubungkan tabel dimensi dengan tabel fakta yang telah dibuat sehingga terbentuk suatu skema *data warehouse*.

Penelitian ini menggunakan *star schema* untuk membentuk *data warehouse* yang akan digunakan. Gambar 27 menunjukkan bentuk *star schema* yang dibuat untuk informasi kecepatan *upload*, dimana tabel *fakta_upload* terhubung dengan tabel *dim_kelurahan* melalui atribut *Kode_Kelurahan* dan tabel *dim_kecamatan* melalui atribut *Kode_Kecamatan*.



Gambar 27. *Star schema* untuk *data warehouse* data kecepatan *upload*

Perancangan *star schema* untuk data kecepatan *download* hampir sama dengan *star schema* pada data kecepatan *upload*. Perbedaan diantara kedua *data warehouse* tersebut terletak pada tabel fakta yang digunakan, dimana pada data kecepatan *download* digunakan tabel *fakta_download*. Gambar 28 menunjukkan *star schema* pada data kecepatan *download*.



Gambar 28. *Star schema* untuk *data warehouse* data kecepatan *download*

Penelitian ini menggunakan MySQL dalam merancang *database* sebagai tempat penyimpanan tabel dimensi dan tabel fakta yang akan digunakan. Tabel dimensi dan tabel fakta merupakan tabel yang akan membentuk skema *data warehouse* yang akan dibuat. Untuk membuat *database*, dibutuhkan beberapa aplikasi pendukung yaitu XAMPP dan MySQL.

Penelitian ini akan menggunakan *star schema data warehouse* untuk proses pembuatan *data warehouse*. Pembuatan *data warehouse* akan dilakukan setelah proses *ETL* selesai. Proses diawali dengan penentuan *Primary Key (PK)* pada tabel dimensi dan *Foreign Key (FK)* pada tabel fakta. Penelitian ini menggunakan 2 *data warehouse* yaitu *data warehouse* untuk data kecepatan *upload* dan *data warehouse* untuk kecepatan *download*. Gambar 28 menunjukkan skema *data warehouse* yang akan dibuat pada penelitian ini.

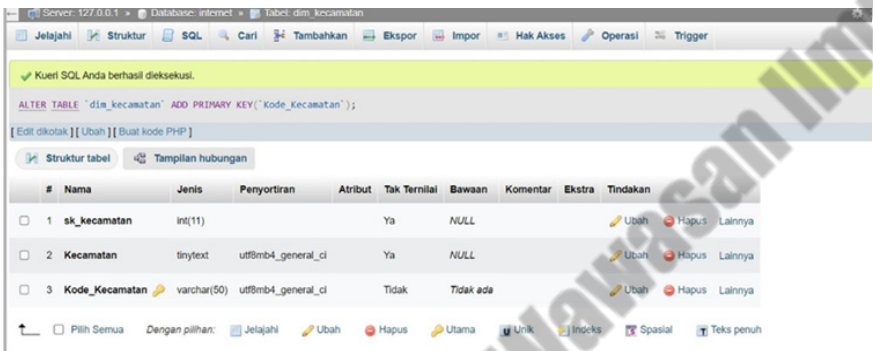
Langkah awal dalam penyusunan *data warehouse* yang dibutuhkan yaitu melakukan deklarasi *grain* yang akan menjadi dasar dalam menentukan tabel dimensi dan tabel fakta yang akan dibuat. Tabel 11 menunjukkan pemilihan *grain* terkait dengan penentuan informasi yang akan dimunculkan pada tabel fakta. Pada tabel tersebut ditentukan tabel fakta yang akan dibuat adalah informasi tentang kecepatan internet *upload* dan *download* dari setiap *provider*. *Data warehouse* yang akan dibuat pada penelitian ini menggunakan *star schema* yang terdiri dari tabel dimensi *dim_* kecamatan dan *dim* kelurahan, serta tabel fakta untuk kecepatan *upload* dan *download*.

Tabel 17. Tabel grain dalam pembentukan data warehouse

Dimensi	Grain	
	Kecepatan Upload	Kecepatan Download
Kecamatan	V	V
Kelurahan	V	V

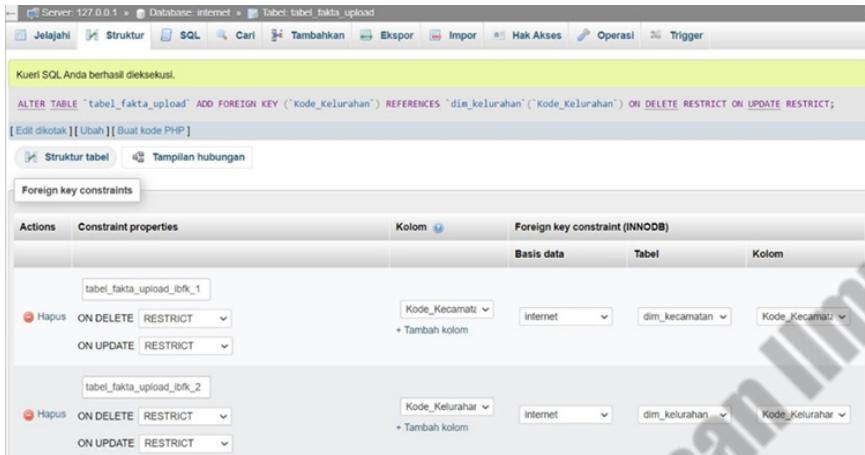
Tabel dimensi dan tabel fakta yang akan digunakan selanjutnya digabungkan menjadi suatu skema yang disebut *star schema*.

Pembentukan diawali dengan penentuan atribut yang menjadi *primary key* pada tabel dimensi dan *foreign key* pada tabel fakta. Gambar 29 menunjukkan cara penentuan *primary key* pada tabel dimensi yang akan digunakan. Pada tabel *dim_kecamatan*, atribut *Kode_Kecamatan* yang akan dijadikan *primary key* dipilih dan dijadikan utama seperti pada gambar. Hal serupa juga dilakukan untuk tabel *dim_kelurahan*.



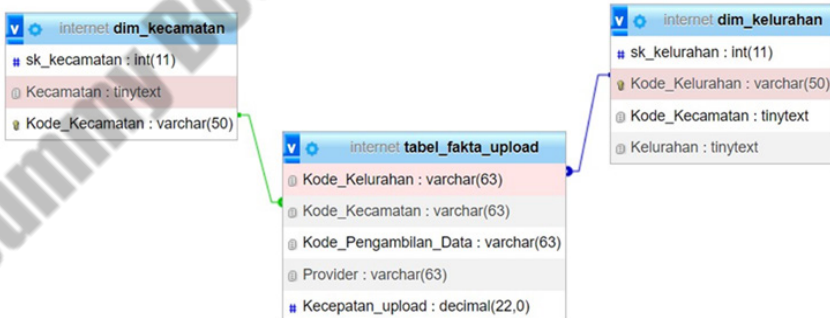
Gambar 29. Proses penentuan primary key pada tabel dimensi

Penentuan *foreign key* pada tabel fakta dilakukan untuk menghubungkan atribut yang sama pada tabel dimensi dan tabel fakta. Proses diawali dengan memilih tabel fakta upload dan pilih struktur. Langkah selanjutnya adalah menuliskan *query* yang berfungsi untuk menentukan atribut yang akan dijadikan *foreign key*, yaitu tabel *Kode_Kecamatan* dan *Kode_Kelurahan*. Gambar 30 berikut ini menunjukkan *query* yang dituliskan dan hasil yang didapatkan untuk tabel *fakta_upload*.



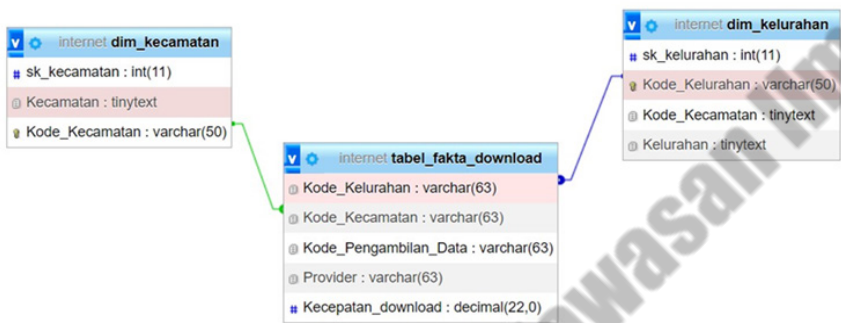
Gambar 30. Proses penentuan *foreign key*

Penentuan *primary key* dan *foreign key* pada proses sebelumnya digunakan untuk menghubungkan tabel dimensi dengan tabel fakta yang telah dibuat sehingga terbentuk suatu skema *data warehouse*. Penelitian ini menggunakan *star schema* untuk membentuk *data warehouse* yang akan digunakan. Gambar 31 menunjukkan bentuk *star schema* yang dibuat untuk informasi kecepatan *upload*, dimana *tabel_fakta_upload* terhubung dengan *dim_kelurahan* melalui atribut *Kode_Kelurahan* dan *dim_kecamatan* melalui atribut *Kode_Kecamatan*.



Gambar 31. Star schema untuk data warehouse data kecepatan upload

Perancangan *star schema* untuk data kecepatan *download* hampir sama dengan *star schema* pada data kecepatan *upload*. Perbedaan diantara kedua *data warehouse* tersebut terletak pada tabel fakta yang digunakan, dimana pada data kecepatan *download* digunakan tabel_fakta_download. Gambar 32 menunjukkan *star schema* pada data kecepatan *download*.



Gambar 32. Star schema untuk data warehouse data kecepatan download

Dummy Book Finish Wawasan Ilmu

BAB IV

PERANCANGAN SISTEM INTELIJENSIA BISNIS

A. Perancangan Sistem Intelijensia Bisnis di Pendidikan Tinggi

Sistem intelijensia bisnis yang disusun terdiri dari serangkaian komponen yang menghasilkan informasi untuk menilai performa publikasi ilmiah melalui visualisasi data, seperti dashboard. Pemanfaatan sistem business intelligence memiliki keunggulan dalam mendukung pengambilan keputusan manajemen secara efisien, khususnya terkait kebijakan pengembangan dosen dalam aspek penyelenggaraan pendidikan, pelatihan, penugasan, promosi jabatan, dan pemberian penghargaan berdasarkan:

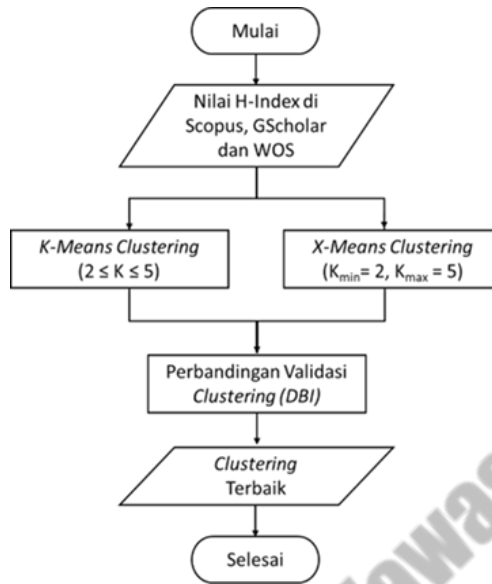
- Perangkingan penulis berdasarkan prestasi mereka dalam menerbitkan karya ilmiah yang mengacu pada Indikator Kinerja Penulis yang tercantum dalam Pedoman Publikasi Ilmiah 2019 Kementerian Riset, Teknologi, dan Pendidikan Tinggi (Sinta).
- Penilaian pencapaian target sitasi publikasi dosen yang sesuai dengan program studi berdasarkan Standar Hasil Penelitian di lingkungan Universitas.

- Penilaian pencapaian target publikasi internasional yang disesuaikan dengan program studi berdasarkan Standar Hasil Penelitian di lingkungan Universitas.
- Evaluasi pencapaian target subjek penelitian dari para penulis yang sejalan dengan rencana pengembangan penelitian, mengacu pada Standar Hasil Penelitian di lingkungan Universitas.

1. Penerapan Clustering

Penerapan *data mining* bertujuan untuk mengelompokkan peneliti berdasarkan peringkat H-Index mereka di platform Scopus, Google Scholar, dan Web of Science. Proses pengelompokkan peneliti dilakukan dengan menggunakan dua algoritma *clustering*, yaitu K-Means dan X-Means. Sumber data yang digunakan untuk mengelompokkan berasal dari *data warehouse* yang telah melalui proses transformasi dan telah disaring.

Dalam proses pemodelan *clustering*, digunakan algoritma K-Means dan X-Means dengan jumlah kelompok (K) sebanyak 5 *cluster*. Pengelompokan dilakukan secara berurutan, dimulai dari nilai K sebesar 2, 3, 4, hingga 5. Proses *clustering* menggunakan perangkat lunak *data mining* khusus. Penelitian ini memanfaatkan program Python sebagai perangkat lunak *data mining*. Detail alur proses *clustering* dapat ditemukan dalam Gambar 33.



Gambar 33. Alur Proses Clustering

Pengujian kelompok menggunakan *Davies Bouldin Index* dilakukan pada K-Means dan X-Means dengan pengelompokan sebanyak 5 kelompok yang dihitung selama proses pengelompokan menggunakan perangkat lunak Python untuk analisis data. Hasil evaluasi kelompok dibandingkan untuk menentukan pengelompokan terbaik. Tabel 12 menunjukkan perbandingan nilai *Davies Bouldin Index* setelah proses pengelompokan.

Tabel 18. Perbandingan Nilai Davies Bouldin Index

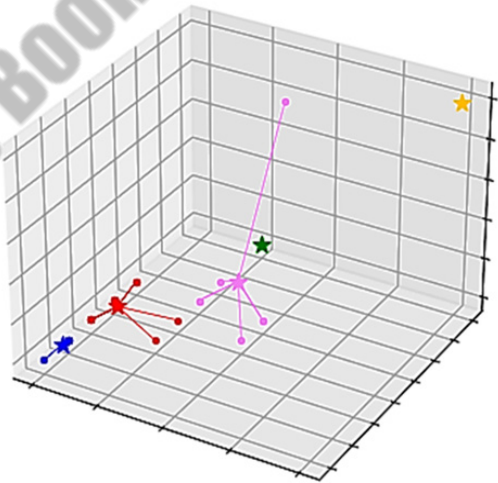
K-Means Clustering		X-Means Clustering		
K/Cluster	Nilai DBI	K	Cluster	Nilai DBI
K = 2	0,778144	K _{min} = 2, K _{max} = 5	4	0,614357
K = 3	0,682411	K _{min} = 3, K _{max} = 5	5	0,537040
K = 4	0,655393	K _{min} = 4, K _{max} = 5	5	0,553758
K = 5	0,734986	K _{min} = 5, K _{max} = 5	5	0,747751

Hasil analisis *Davies Bouldin Index* menunjukkan bahwa algoritma *clustering* X-Means memiliki kinerja pengelompokan terbaik dengan menggunakan 5 cluster ($K_{min} = 3$, $K_{max} = 5$), serta memiliki nilai *Davies Bouldin Index* sekitar 0,537040.

Tabel 19. Hasil Clustering dan Centroid

Cluster	Centroid H-Index Scopus	Centroid H-Index GScholar	Centroid H-Index WOS	Jumlah Peneliti
0	0,125000	3,812500	0,000000	16
1	0,000000	1,777778	0,000000	9
2	1,000000	8,000000	0,000000	1
3	4,000000	8,000000	1,000000	1
4	1,571429	5,142857	0,142857	7

Berdasarkan penggunaan metode pengelompokan dengan menggunakan algoritma X-Means, diperoleh lima kelompok peneliti sebagai hasil terbaik. Informasi tentang hasil pengelompokan peneliti beserta nilai tengah (*centroid*) dari setiap kelompok dapat ditemukan di Tabel 14.



Gambar 34. Visualisasi Pengelompokan Menggunakan X-Means

Peta visualisasi hasil pengelompokan data menggunakan metode X-Means dalam penelitian ini terdapat di Gambar 34.

Tabel 20. Tingkat Pencapaian H-Index

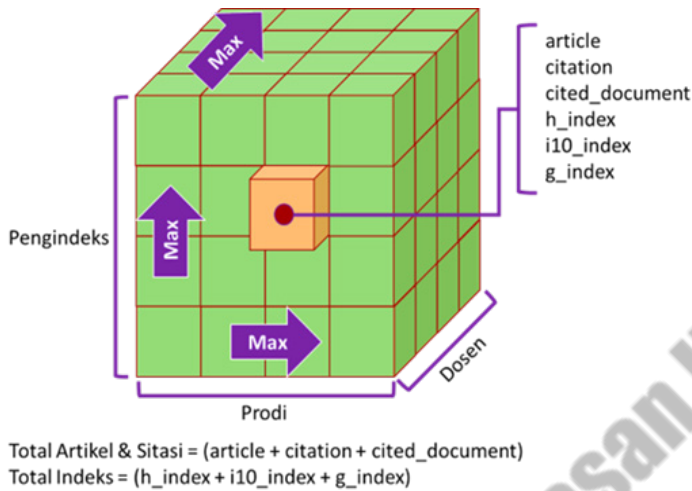
<i>Cluster</i>	<i>H-Index Scopus</i>	<i>H-Index GScholar</i>	<i>H-Index WOS</i>
0	Rendah	Rendah	Terrendah
1	Terrendah	Terrendah	Terrendah
2	Rendah	Tertinggi	Terrendah
3	Tertinggi	Tertinggi	Tertinggi
4	Rendah	Tinggi	Rendah

Berdasarkan nilai titik pusat (*centroid*) di setiap kelompok, dapat disimpulkan bahwa setiap kelompok memiliki tingkat pencapaian H-Index yang berbeda. Informasi mengenai tingkat pencapaian H-Index di setiap kelompok dapat ditemukan dalam Tabel 14.

2. Penerapan OLAP

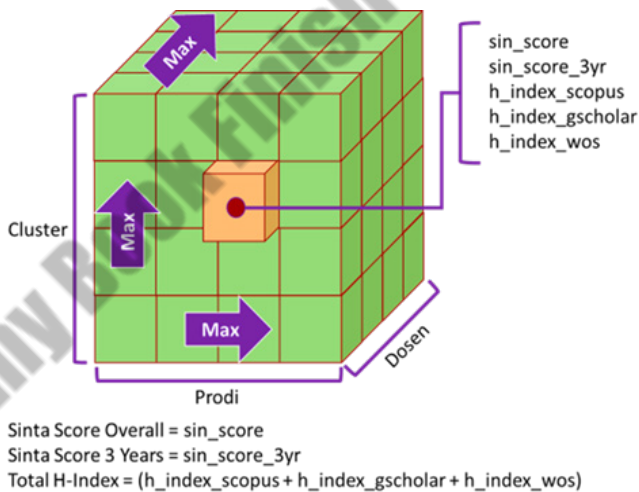
Setelah data telah disimpan dengan baik di dalam *data warehouse*, data tersebut dimanfaatkan untuk mendukung proses pengambilan keputusan serta memberikan solusi atas pertanyaan-pertanyaan bisnis dan manajemen menggunakan teknologi OLAP sebagai metode analisis data dengan melakukan eksekusi *query* analitik multidimensi. Dalam sistem inteligensi bisnis ini, dibentuklah kubus OLAP sebagai teknik analisis data untuk memenuhi kebutuhan fungsional dalam menyediakan informasi yang diperlukan, yaitu:

- *Cube* indeks peneliti digunakan untuk menyajikan data mengenai indeks peneliti dapat dilihat pada Gambar 35.



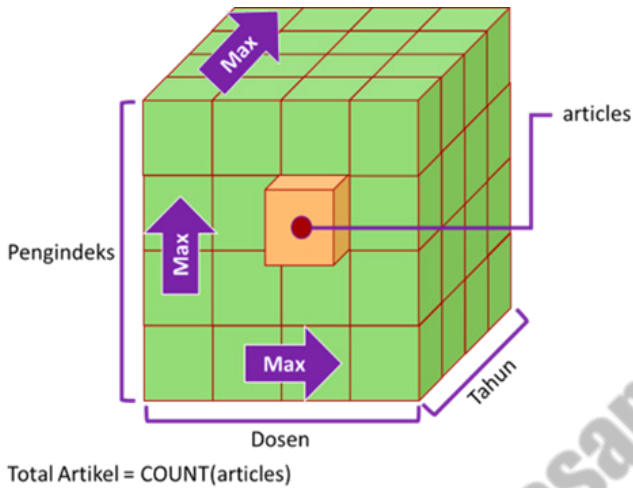
Gambar 35. Cube Indeks Peneliti

- *Cube score* peneliti digunakan untuk menyediakan data tentang *score* peneliti dapat dilihat pada Gambar 36.



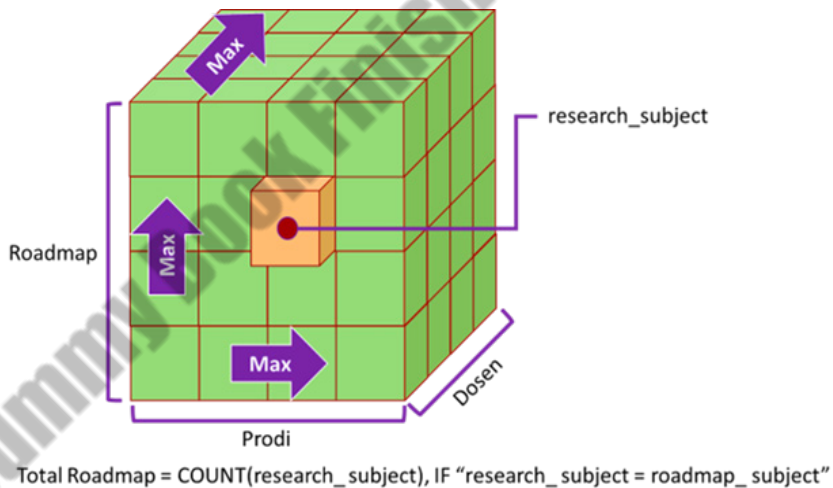
Gambar 36. Cube Score Peneliti

- *Cube* artikel publikasi bertujuan untuk memberikan informasi mengenai artikel yang dipublikasikan dapat dilihat pada Gambar 37.



Gambar 37. Cube Artikel Publikasi

- *Cube* subjek penelitian bertujuan untuk memberikan informasi mengenai subjek penelitian dapat dilihat pada Gambar 38.



Gambar 38. Cube Subjek Penelitian

Setelah pembentukan kubus OLAP selesai, langkah selanjutnya adalah mengambil data untuk keperluan analisis dan menampilkan hasil analisis tersebut pada Sistem Intelijenisa Bisnis.

3. Visualisasi Data

Power BI digunakan sebagai aplikasi perangkat lunak dalam membangun sistem intelijensia bisnis. Penerapan Power BI memungkinkan pengolahan data yang mendetail serta penyajian visual grafis yang interaktif. Tampilan utama dari sistem intelijensia bisnis yang telah dibangun dapat diakses pada Gambar 39.



Gambar 39. Tampilan Halaman Utama Sistem Intelijensia Bisnis

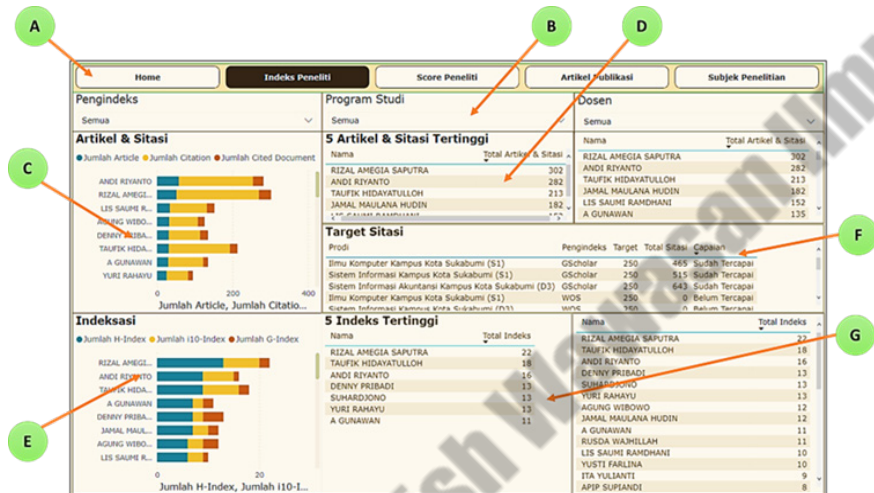
- Informasi Indeks Peneliti

Pada halaman Indeks Peneliti terdapat komponen-komponen visualisasi data, yaitu:

- Page Navigator*, sebagai alat navigasi dari suatu halaman berpindah ke halaman lain di dalam visualisasi data.
- Slicer*, sebagai alat untuk menyaring data berdasarkan pengindeks, program studi, dan dosen.
- Grafik Artikel dan Sitasi, menampilkan visualisasi data jumlah artikel, jumlah sitasi, jumlah artikel tersitasi.
- Hasil analisis 5 nilai total artikel dan sitasi tertinggi.
- Grafik Indeksasi, menampilkan visualisasi data jumlah *H-Index*, jumlah *i10-Index*, jumlah *G-Index*.

- f. Capaian target sitasi publikasi dosen sesuai program studi.
- g. Hasil analisis 5 nilai total indeks tertinggi.

Tampilan halaman Indeks Peneliti pada sistem *business intelligence* yang dikembangkan dapat dilihat pada Gambar 40.



Gambar 40. Tampilan Halaman Indeks Peneliti

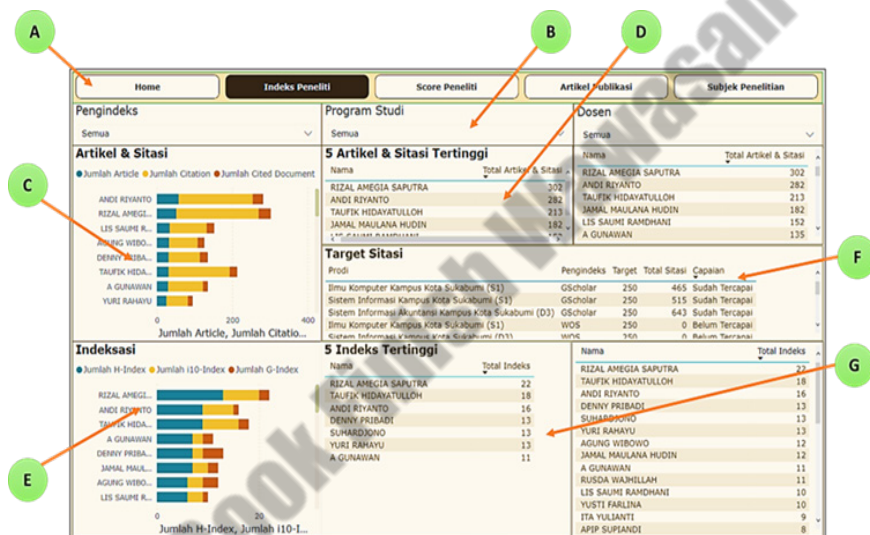
- Informasi Score Peneliti

Pada halaman *Score Peneliti* terdapat komponen-komponen visualisasi data, yaitu:

- a. *Page Navigator*, sebagai alat navigasi dari suatu halaman berpindah ke halaman lain di dalam visualisasi data.
- b. *Slicer*, sebagai alat untuk menyaring data berdasarkan program studi, dan dosen.
- c. Grafik *Sinta Score*, menampilkan visualisasi data nilai *sinta score* dan *sinta score 3yr* untuk masing-masing peneliti.
- d. Hasil analisis 5 nilai *sinta score* tertinggi.
- e. Hasil analisis 5 nilai *sinta score 3yr* tertinggi.

- f. Grafik *H-Index*, menampilkan visualisasi data nilai *H-Index* dari setiap pengindeks (Scopus, G-Scholar, WOS) untuk masing-masing peneliti.
- g. Hasil analisis 5 nilai total *H-Index* tertinggi.
- j. Grafik Elemen *Cluster*, menampilkan visualisasi data jumlah dan persentase anggota dari setiap *cluster* peneliti hasil pengelompokan menggunakan *data mining*.

Tampilan halaman *Score Peneliti* pada sistem *business intelligence*.



Gambar 41. Tampilan Halaman *Score Peneliti*

- Informasi Artikel Publikasi

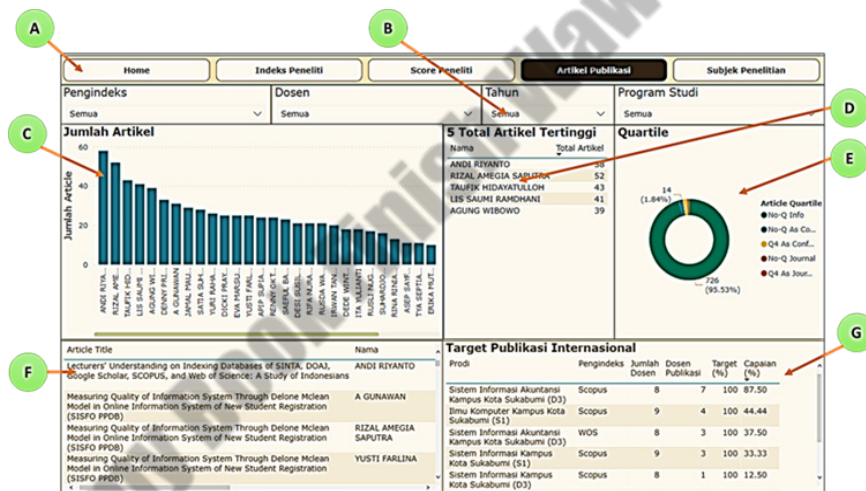
Pada halaman Artikel Publikasi terdapat komponen-komponen visualisasi data, yaitu:

- a. *Page Navigator*, sebagai alat navigasi dari suatu halaman berpindah ke halaman lain di dalam visualisasi data.
- b. *Slicer*, sebagai alat untuk menyaring data berdasarkan pengindeks, tahun, program studi, dan dosen.
- c. Grafik Jumlah Artikel, menampilkan visualisasi data jumlah

artikel yang terindeks dari setiap pengindeks (Scopus, G-Scholar, WOS).

- d. Hasil analisis 5 nilai total artikel tertinggi.
- e. Grafik *Quartile*, menampilkan visualisasi data jumlah dan persentase *quartile* artikel yang tercatat pada pengindeks.
- f. Tabel Judul Artikel, menampilkan daftar judul artikel yang dapat disaring berdasarkan pengindeks, tahun, program studi, dosen, dan *quartile*.
- g. Capaian target publikasi internasional sesuai program studi.

Tampilan halaman Artikel Publikasi pada sistem *business intelligence* yang dikembangkan dapat dilihat pada Gambar 42.



Gambar 42. Tampilan Halaman Artikel Publikasi

- Informasi Subjek Penelitian

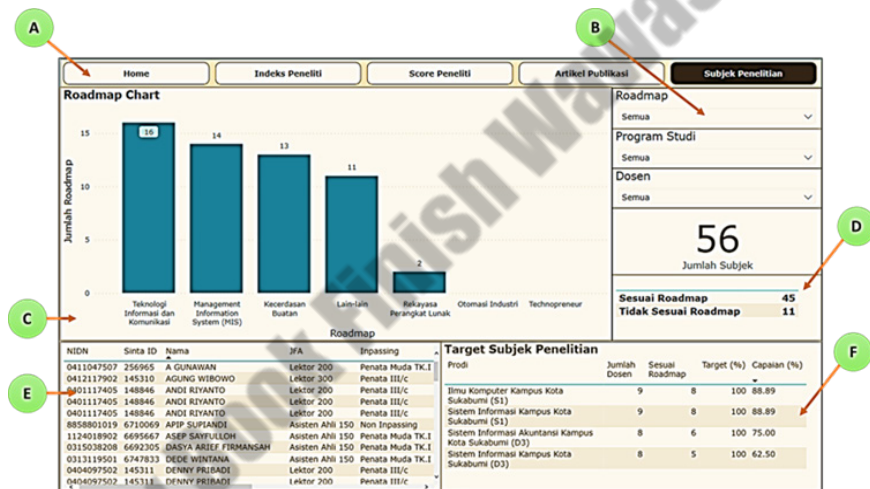
Pada halaman Subjek Penelitian terdapat komponen-komponen visualisasi data, yaitu:

- a. *Page Navigator*, sebagai alat navigasi dari suatu halaman berpindah ke halaman lain di dalam visualisasi data.
- b. *Slicer*, sebagai alat untuk menyaring data berdasarkan *roadmap*,

program studi, dan dosen.

- Grafik *Roadmap*, menampilkan visualisasi data jumlah peneliti berdasarkan *roadmap* penelitian.
- Hasil analisis kesesuaian subjek dengan roadmap penelitian.
- Tabel Peneliti, menampilkan daftar peneliti yang dapat disaring berdasarkan *roadmap*, program studi, dan dosen.
- Capaian target subjek penelitian dosen yang sesuai dengan roadmap penelitian.

Tampilan halaman Subjek Penelitian pada sistem *business intelligence* yang dikembangkan dapat dilihat pada Gambar 43.



Gambar 43. Tampilan Halaman Subjek Penelitian

4. Evaluasi Sistem

Uji verifikasi dan uji validasi dilakukan pada evaluasi sistem intelijensia bisnis. Evaluasi sistem bertujuan untuk membuktikan bahwa tidak ada kesalahan pada sistem dan sebagai pembuktian bahwa sistem dapat memberikan output sesuai dengan harapan dan kebutuhan pengguna. Uji verifikasi sebagai pembuktian bahwa tidak ada kesalahan pada sistem dan berjalan sesuai spesifikasi yang

telah ditentukan. Berdasarkan hasil uji verifikasi, sistem intelijensia bisnis dinyatakan berhasil diverifikasi dan terbukti tidak ada kesalahan pada sistem. Hasil uji verifikasi sistem dapat dilihat pada Tabel 15.

Tabel 21. Hasil Uji Verifikasi Sistem

Daftar Uji	Langkah Pengujian	Indikator Keberhasilan	Hasil Pengujian	Simpulan
Menguji proses <i>web scraping</i> .	Menjalankan program <i>web scraping</i> .	- Tidak ada <i>bug</i> saat menjalankan program. - Hasil <i>web scraping</i> tersimpan dalam <i>database</i>.	Tidak ada <i>bug</i> saat proses <i>web scraping</i> dan data tersimpan dalam <i>database</i> .	Berhasil
Menguji proses <i>ETL</i> dan <i>data warehouse</i> .	Menjalankan program <i>ETL</i> .	- Tidak ada <i>bug</i> saat menjalankan program. - Hasil <i>ETL</i> tersimpan dalam <i>data warehouse</i>.	Tidak ada <i>bug</i> saat proses <i>ETL</i> dan data tersimpan dalam <i>data warehouse</i> .	Berhasil
Menguji proses <i>clustering</i> .	Menjalankan program <i>clustering</i> .	- Tidak ada <i>bug</i> saat menjalankan program. - Hasil <i>clustering</i> tersimpan dalam <i>data warehouse</i>.	Tidak ada <i>bug</i> saat proses <i>clustering</i> dan data tersimpan dalam <i>data warehouse</i> .	Berhasil

Menguji <i>dashboard</i> Indeks Peneliti.	Menampilkan <i>dashboard</i> Indeks Peneliti.	- Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i>. - Informasi indeks peneliti tampil dengan baik.	Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i> dan informasi indeks peneliti tampil dengan baik.	Berhasil
Menguji <i>dashboard</i> Score Peneliti.	Menampilkan <i>dashboard</i> Score Peneliti.	- Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i>. - Informasi <i>score</i> peneliti tampil dengan baik.	Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i> dan informasi <i>score</i> peneliti tampil dengan baik.	Berhasil
Menguji <i>dashboard</i> Artikel Publikasi	Menampilkan <i>dashboard</i> Artikel Publikasi	- Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i>. - Informasi artikel publikasi tampil dengan baik.	Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i> dan informasi artikel publikasi tampil dengan baik.	Berhasil
Menguji <i>dashboard</i> Subjek Penelitian.	Menampilkan <i>dashboard</i> Subjek Penelitian.	- Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i>. - Informasi subjek penelitian tampil dengan baik.	Tidak ada <i>bug</i> saat menampilkan <i>dashboard</i> dan informasi subjek penelitian tampil dengan baik.	Berhasil

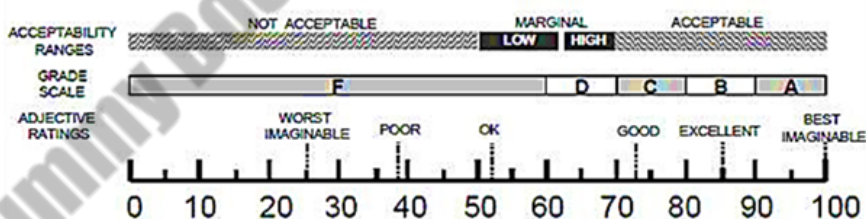
Uji validasi rancangan sistem dilakukan untuk membuktikan bahwa sistem mampu menyajikan informasi sesuai dengan harapan dan kebutuhan pengguna menggunakan metode *System Usability Scale (SUS)*. Metode *SUS* terdiri dari 10 pertanyaan dan 5 pilihan jawaban dengan bobot nilai antara 0 sampai 4 menggunakan skala Likert. Pertanyaan dengan nomor ganjil memiliki makna positif, sedangkan pertanyaan dengan nomor genap memiliki makna negatif (Atsani et al., 2019). Daftar pertanyaan kuesioner *SUS* dapat dilihat pada Tabel 4.5.

Tabel 22. Daftar Pertanyaan Kuesioner SUS

No	Pertanyaan	Jawaban				
		STS	TS	RG	ST	SS
		1	2	3	4	5
1	Saya berpikir akan menggunakan sistem ini lagi.					
2	Saya merasa sistem ini rumit untuk digunakan.					
3	Saya merasa sistem ini mudah digunakan.					
4	Saya membutuhkan bantuan dari orang lain atau teknisi dalam menggunakan sistem ini.					
5	Saya merasa fitur-fitur sistem ini berjalan dengan semestinya.					
6	Saya merasa ada banyak hal yang tidak konsisten (tidak serasi pada sistem ini).					
7	Saya merasa orang lain akan memahami cara menggunakan sistem ini dengan cepat.					

8	Saya merasa sistem ini membingungkan.					
9	Saya merasa tidak ada hambatan dalam menggunakan sistem ini.					
10	Saya perlu membiasakan diri terlebih dahulu sebelum menggunakan sistem ini.					
STS : Sangat Tidak Setuju, TS : Tidak Setuju RG : Ragu-ragu, ST : Setuju, SS : Sangat Setuju						

Uji verifikasi dan uji validasi dilakukan pada evaluasi sistem diisi oleh ketua program studi dan ketua LPPM. Nilai pada pertanyaan bernomor ganjil dihitung dengan cara nilai jawaban dikurangi dengan nilai 1, sedangkan nilai pada pertanyaan bernomor genap dihitung dengan cara 5 dikurangi nilai jawaban. Total nilai akhir dihitung dengan cara menjumlahkan total nilai setiap pertanyaan. Total nilai akhir dikalikan dengan 2,5 dan dibagi dengan jumlah responden untuk mendapatkan nilai *SUS* antara 0 sampai 100. Sistem dikategorikan baik jika nilai akhir *SUS* ≥ 70 (Atsani et al., 2019). Grafik peringkat persentil nilai *SUS* dapat dilihat pada Gambar 44.



Gambar 44. Grafik Peringkat Persentil Nilai SUS

Jawaban kuesioner dari setiap responden bisa dilihat pada Tabel 4.6 dan hasil perhitungan nilai *SUS* dapat dilihat pada Tabel 17.

Tabel 23. Jawaban Kuesioner SUS

Responden	P1	P2	P3	P4	P5	P6	P7	P8	P9	P10
Kaprodi Ilmu Komputer (S1)	5	2	4	1	4	1	5	2	5	1
Kaprodi Sistem Informasi Akuntansi (D3)	5	2	4	2	5	1	4	1	4	1
Kaprodi Sistem Informasi (D3)	5	2	4	1	5	2	4	2	5	1
Kaprodi Sistem Informasi (S1)	4	1	4	1	4	2	4	2	4	2
Ketua LPPM	5	1	4	1	4	1	5	2	4	2

Tabel 24. Perhitungan Nilai SUS

Responden	P1	P2	P3	P4	P5	P6	P7	P8	P9	P10	Total
Kaprodi Ilmu Komputer (S1)	4	3	3	4	3	4	4	3	4	4	36
Kaprodi Sistem Informasi Akuntansi (D3)	4	3	3	3	4	4	3	4	3	4	35

Kaprodi Sistem Informasi (D3)	4	3	3	4	4	3	3	3	4	4	35
Kaprodi Sistem Informasi (S1)	3	4	3	4	3	3	3	3	3	3	32
Ketua LPPM	4	4	3	4	3	4	4	3	3	3	35
Total Nilai Akhir											173
Nilai SUS	$(173 \times 2,5) / 5$										86,5

Hasil perhitungan nilai *SUS* dibandingkan dengan grafik peringkat persentil nilai *SUS* dapat disimpulkan bahwa sistem intelijensia bisnis dengan nilai 86,5 termasuk dalam kategori baik sekali serta layak digunakan.

Model Sistem Business Intelligence Publikasi Ilmiah (SiBIPI) berhasil dikembangkan dengan menggunakan *data warehouse* yang memanfaatkan model dimensional. Model ini mencakup model indeks peneliti, model score peneliti, model artikel publikasi, dan model subjek penelitian.

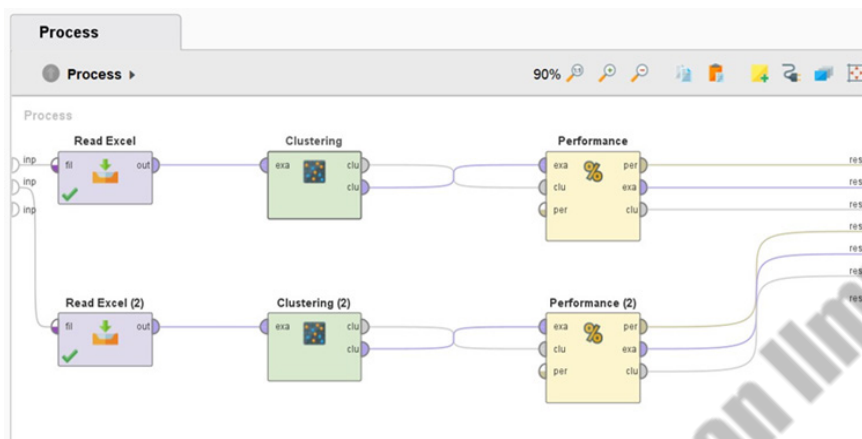
Pengukuran besar data dan waktu pemrosesan data menunjukkan bahwa model *star schema*, dengan ukuran data 336 KB dan waktu proses 0,00554 detik, lebih efisien dibandingkan dengan model *snowflakes schema* yang memiliki ukuran data 368 KB dan waktu proses 0,00611 detik. Pengukuran validasi *cluster* menggunakan Davies Bouldin Index menunjukkan bahwa algoritma *X-Means clustering* menghasilkan kinerja *clustering* terbaik dengan 5 *cluster* optimal ($K_{min}=3$, $K_{max}=5$) dan nilai Davies Bouldin Index sebesar 0,537040. Hasil pengelompokan menunjukkan bahwa *cluster* 0 terdiri dari 6 orang, *cluster* 1 terdiri dari 9 orang, *cluster* 2 terdiri dari 1 orang, *cluster* 3 terdiri dari 1 orang, dan *cluster* 4 terdiri dari 7 orang.

Implementasi OLAP menghasilkan empat *cube*, yaitu *cube* indeks peneliti, *cube* score peneliti, *cube* artikel publikasi, dan *cube* subjek penelitian. Model sistem *business intelligence* ini bermanfaat untuk mengukur kinerja publikasi ilmiah sesuai dengan standar hasil penelitian, yang dapat menjadi dasar dalam pengambilan keputusan optimal untuk pengembangan dosen. Prototipe sistem *business intelligence* berhasil dibuat dengan tampilan *dashboard* menggunakan Power BI. Hasil uji verifikasi menunjukkan bahwa sistem ini tidak memiliki kesalahan, dan hasil perhitungan nilai *System Usability Scale* (SUS) pada uji validasi adalah 86,5, yang menunjukkan bahwa sistem *business intelligence* ini termasuk dalam kategori *excellent* dan layak digunakan.

B. Perancangan Sistem Intelijensia Bisnis di Industri Provider

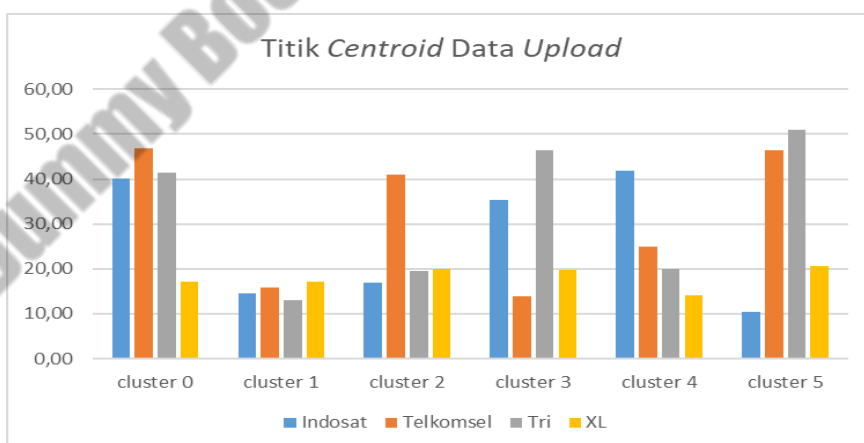
1. Proses *Clustering*

Proses *clustering* pada penelitian ini dilakukan menggunakan aplikasi *Rapidminer* menggunakan data hasil proses OLAP baik data *upload* maupun data *download*. Proses *clustering* diawali dengan *import* data yang akan di *cluster* ke dalam operator *Read Excel*. Data tersebut kemudian diolah melalui operator *Clustering*. Operator *Clustering* yang digunakan yaitu *K-Means* untuk data *upload* dan *K-Medoids* untuk data *download* sesuai dengan ketentuan yang telah dibuat pada proses penentuan jumlah *cluster*. Gambar 48 menunjukkan proses *clustering* pada penelitian ini menggunakan *Rapidminer*.



Gambar 45. Proses *clustering* pada data *upload* dan data *download*

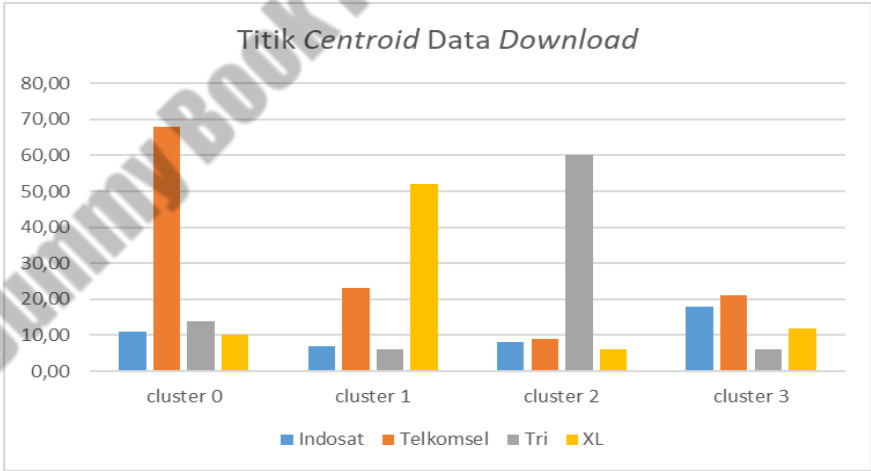
Proses *clustering* yang dilakukan pada penelitian ini bertujuan untuk mengelompokkan wilayah kelurahan kota bekasi berdasarkan data kecepatan *internet*. Hasil *clustering* yang didapatkan dikelompokkan berdasarkan jarak terhadap titik *centroid* pada masing-masing *cluster*. Penentuan titik *centroid* untuk masing-masing atribut yang diteliti menunjukkan karakteristik dari setiap atribut pada masing-masing *cluster*. Gambar 49 dan gambar 50 menunjukkan titik *centroid* masing-masing *cluster* untuk data *upload* dan *download*.



Gambar 46. Titik centroid data upload

Gambar 47 diatas menunjukkan karakteristik masing-masing *cluster* untuk data *upload*. Pada *cluster* 1, P1, P2, dan P3 memiliki titik *centroid* yang tinggi sedangkan pada provider P4 memiliki titik *centroid* rendah. *Cluster* 1 memiliki titik *centroid* yang rendah pada semua *provider*. *Cluster* 2 memiliki titik *centroid* yang tinggi pada P2. Titik *centroid* yang tinggi pada *cluster* 3 dimiliki oleh P1 dan P3, sedangkan P2 dan P4 memiliki titik *centroid* yang rendah. *Cluster* 4 memiliki nilai *centroid* yang tinggi hanya pada P1, sedangkan *provider* lainnya memiliki titik *centroid* yang rendah. *Cluster* 5 memiliki titik *centroid* yang tinggi pada *provider* P2 dan P3, sedangkan P1 dan P4 memiliki titik *centroid* yang rendah.

Berdasarkan titik *centroid* yang diperoleh, pengelompokan data dapat dilakukan dengan menentukan jarak terdekat dari setiap data terhadap titik *centroid* yang ada. Data-data tersebut akan dikelompokkan kedalam *cluster* yang memiliki jarak titik *centroid* terdekat dengan data tersebut. Tabel pada lampiran memperlihatkan setiap data kelurahan yang diteliti beserta *cluster* masing-masing kelurahan untuk data kecepatan *internet upload* dan *download*. Data tersebut menjadi salah satu data yang akan dimunculkan pada *dashboard* untuk melakukan visualisasi data pengelompokan kelurahan pada aplikasi *Power BI*.

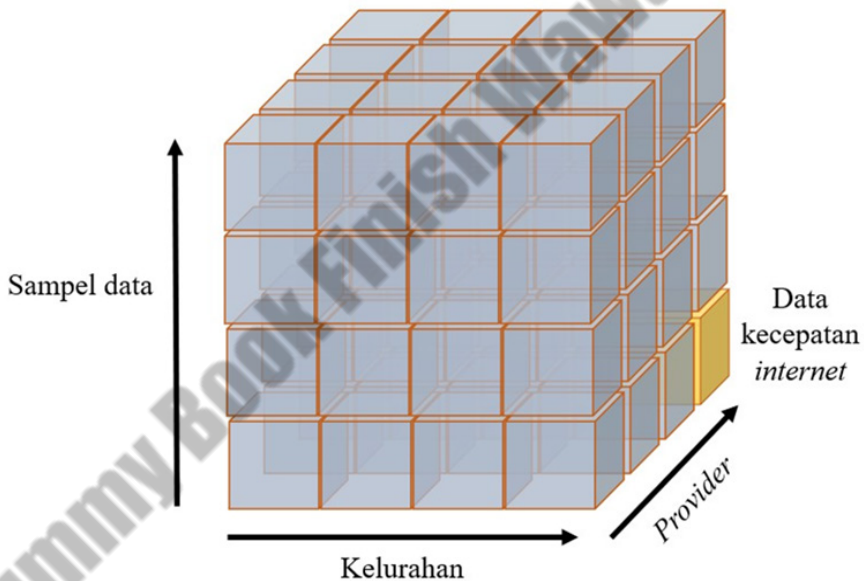


Gambar 47. Titik centroid data download

Gambar 4.16 di atas merupakan gambaran perbandingan titik *centroid* dari masing-masing data *provider* untuk setiap *cluster* yang terbentuk. *Cluster 0* memiliki titik *centroid* yang tinggi pada data P2, dimana *cluster 1* memiliki titik *centroid* yang tinggi hanya pada data P4 dan *cluster 2* memiliki titik *centroid* yang tinggi pada data P3. *Cluster 3* cukup berbeda karena titik *centroid* pada data keempat *provider* bernilai kecil.

2. Penerapan OLAP

Perancangan *OLAP* pada penelitian ini dimulai dengan membuat model *OLAP Cube* dari data yang akan digunakan. Gambar 51 menunjukkan gambaran *OLAP Cube* dengan 3 dimensi yaitu data kelurahan, data *provider*, dan sampel data.



Gambar 48. *OLAP Cube* data kecepatan internet

Perancangan *OLAP* pada penelitian ini dilakukan pada kedua tabel fakta pada masing-masing *data warehouse*. Proses perancangan *OLAP Cube* dilakukan menggunakan *Python* untuk mendapatkan nilai median kecepatan internet *upload* dan *download*. Proses perancangan *OLAP Cube* pada data *upload* menggunakan kode

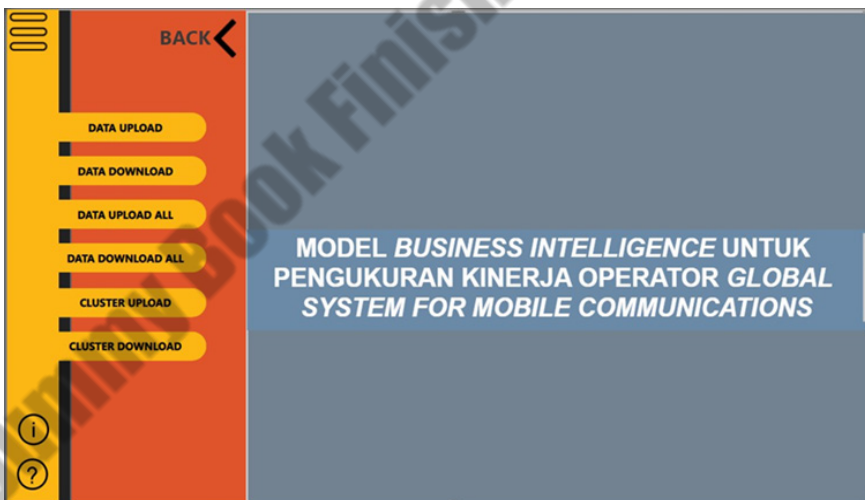
program yang dapat dilihat pada lampiran 3.

Program yang dibuat diawali dengan pemanggilan fungsi untuk koneksi *MySQL* dengan *Python* kemudian memanggil *database internet* dimana tabel_fakta_upload berada. Langkah selanjutnya adalah memanggil atribut pada tabel_fakta_upload yaitu Kode_Kecamatan, Kode_Kelurahan, Kode_Pengambilan_Data, Provider, dan Kecepatan_upload. Langkah terakhir dari proses perancangan *OLAP Cube* untuk tabel_fakta_upload yaitu melakukan fungsi *pivot* sehingga dihasilkan tabel pada lampiran 4.

3. Visualisasi Data

- Halaman Awal dan Overview

Halaman ini dibuat sebagai tampilan awal dan sarana menuju *dashboard* berikutnya. Halaman ini berisi menu menuju halaman awal dan juga menuju halaman Data Upload, Data Download, Data Upload All, Data Download All, Cluster Upload, dan Cluster Download. Gambar 52 merupakan gambaran dari halaman awal yang dibuat.



Gambar 49. Tampilan *dashboard* awal atau *overview*

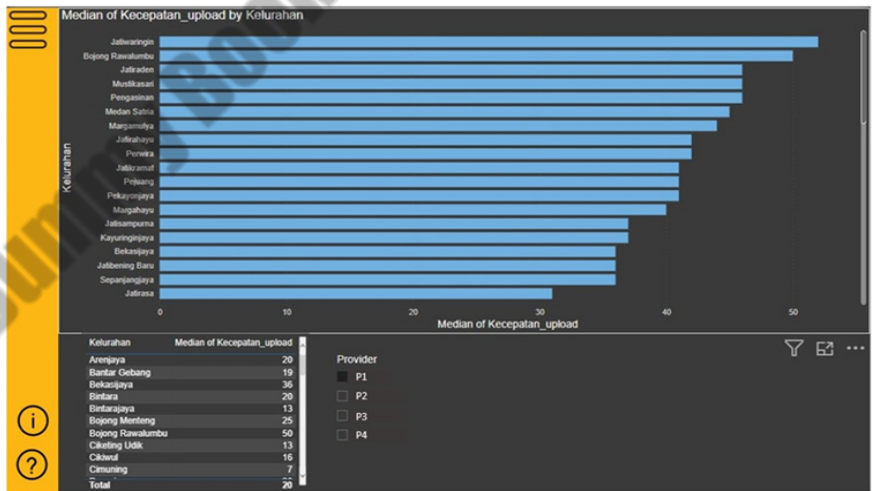
- Halaman Data Upload

Halaman Data Upload menggunakan data pada tabel_fakta_upload yang berisi data kecepatan *internet* saat melakukan *upload*.

Dashboard yang dibuat menampilkan grafik data kecepatan *internet* dan juga tabel. Data yang ditampilkan dapat dipilih berdasarkan *provider* yang akan ditampilkan menggunakan operator *Filter* pada *Power BI*. Dari data yang ditampilkan pada halaman *dashboard* ini, terlihat bahwa kecepatan internet tertinggi untuk *upload* saat menggunakan provider P1 berada di kelurahan Jatiwaringin dengan 52Mbps, sedangkan kelurahan dengan kecepatan *upload* terendah terdapat di kelurahan Telukpucung dengan 1Mbps. Data mengenai kecepatan *upload* tertinggi dan terendah dapat dilihat pada tabel 30.

Tabel 25. Tabel informasi kecepatan upload tertinggi dan terendah

Provider	Status	Kecepatan	Wilayah
P1	Tertinggi	52	Jatiwaringin
	Terendah	1	Telukpucung
P2	Tertinggi	54	Cimuning
	Terendah	3	Jatikarya
P3	Tertinggi	55	Cimuning
	Terendah	3	Telukpucung
P4	Tertinggi	22	Bojongmenteng
	Terendah	3	Jatikramat



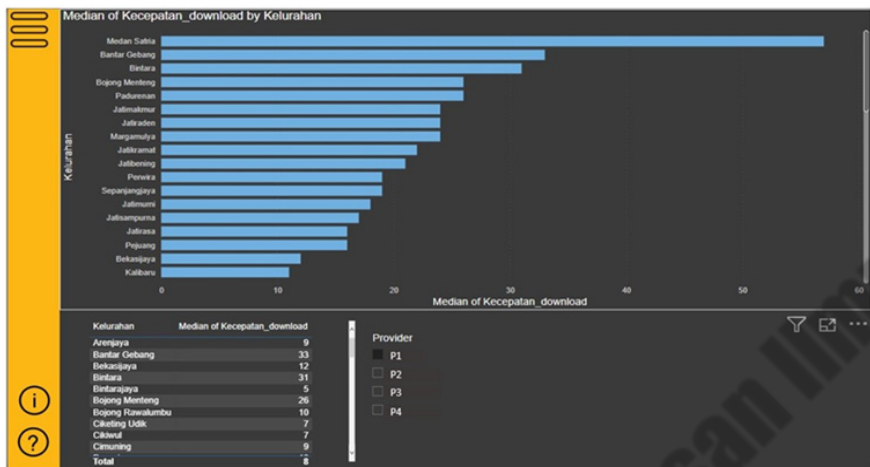
Gambar 50. Tampilan *dashboard* kecepatan *upload*

- Halaman Data Download

Halaman Data *Download* menggunakan data pada tabel_fakta_download yang berisi data kecepatan *internet* saat melakukan *download*. *Dashboard* yang dibuat menampilkan grafik data kecepatan *internet* dan juga tabel. Data yang ditampilkan dapat dipilih berdasarkan *provider* yang akan ditampilkan menggunakan operator *Filter* pada *Power BI*. Dari data yang ditampilkan pada halaman *dashboard* ini, terlihat bahwa kecepatan internet tertinggi untuk *upload* saat menggunakan provider P1 berada di kelurahan Medan Satria dengan 57Mbps, sedangkan kelurahan dengan kecepatan *upload* terendah terdapat di kelurahan Harapanmulya dengan 1Mbps. Data mengenai kecepatan *upload* tertinggi dan terendah dapat dilihat pada tabel 31.

Tabel 26. Tabel informasi kecepatan download tertinggi dan terendah

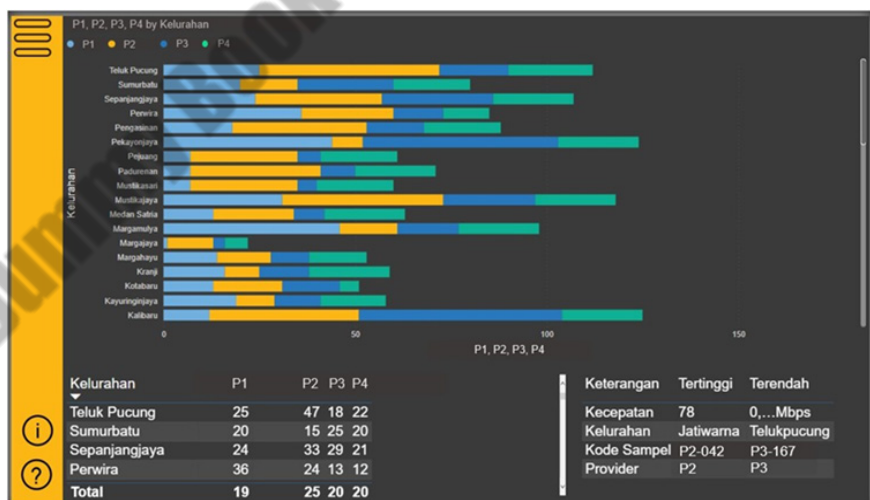
Provider	Status	Kecepatan	Wilayah
P1	Tertinggi	57	Medan Satria
	Terendah	1	Harapanmulya
P2	Tertinggi	68	Mustikajaya
	Terendah	3	Bojong Rawalumbu
P3	Tertinggi	64	Medan Satria
	Terendah	2	Kotabaru
P4	Tertinggi	72	Pengasinan
	Terendah	2	Jatikramat



Gambar 51. Tampilan *dashboard* kecepatan *download*

- Halaman Data *Upload All*

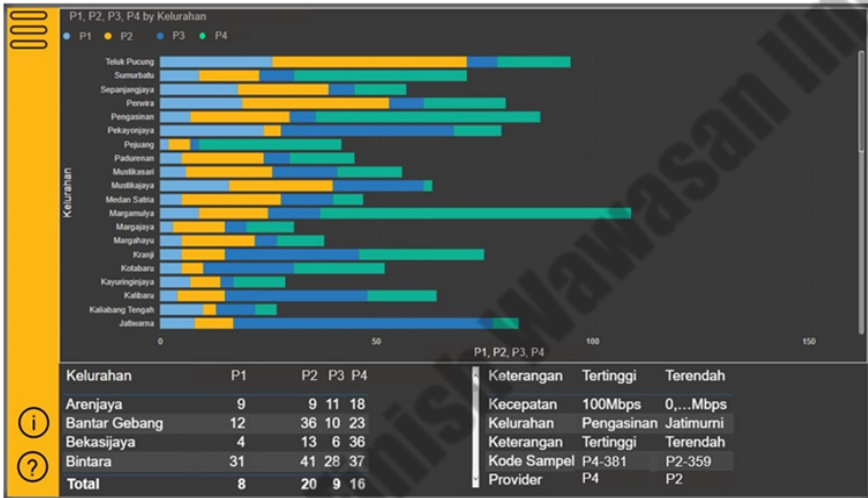
Halaman ini menampilkan data kecepatan *upload* secara keseluruhan untuk keempat *provider* yang diteliti. Data kecepatan *upload* untuk keempat *provider* disatukan dalam satu diagram batang untuk masing-masing kelurahan sehingga bisa terlihat secara visual *provider* yang memiliki kecepatan *upload* tertinggi dan juga terendah.



Gambar 52. Tampilan *dashboard* kecepatan *upload*

- **Halaman Data *Download All***

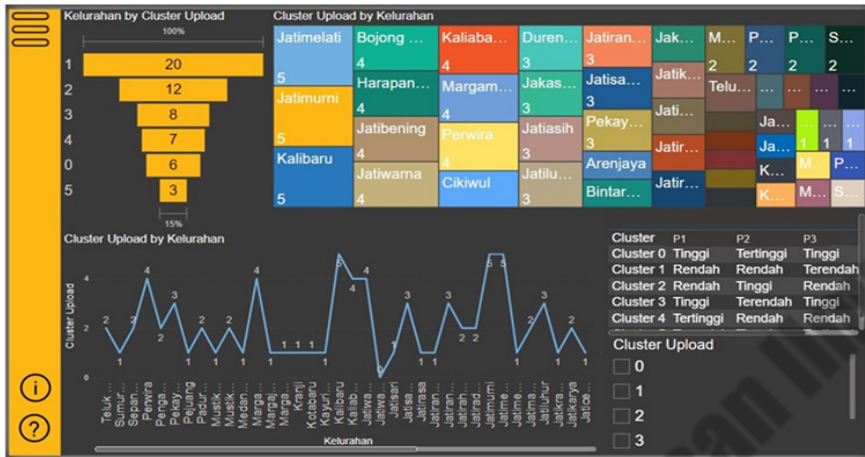
Halaman ini menampilkan data kecepatan *download* secara keseluruhan untuk keempat *provider* yang diteliti. Data kecepatan *download* untuk keempat *provider* disatukan dalam satu diagram batang untuk masing-masing kelurahan sehingga bisa terlihat secara visual *provider* yang memiliki kecepatan *download* tertinggi dan juga terendah.



Gambar 53. Tampilan *dashboard* kecepatan *download*

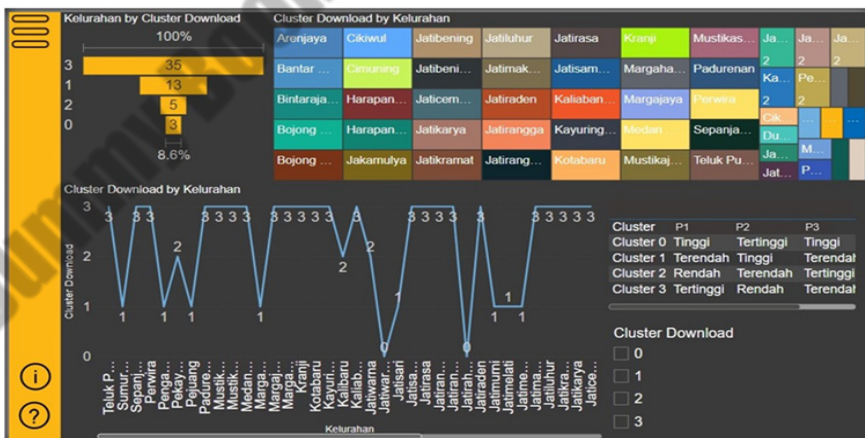
- **Halaman *Cluster Upload***

Gambar 57 di bawah ini memperlihatkan hasil *clustering* yang dilakukan menggunakan data kecepatan *upload*. Dari gambar tersebut terlihat bahwa *cluster* 0 memiliki jumlah anggota sebanyak 6 kelurahan, *cluster* 1 memiliki jumlah anggota sebanyak 20 kelurahan, *cluster* 2 dengan 12 kelurahan, *cluster* 3 dengan 8 kelurahan, *cluster* 4 dengan 7 kelurahan, dan *cluster* 5 dengan 3 kelurahan.



- Halaman *Cluster Download*

Gambar 58 di bawah ini memperlihatkan hasil *clustering* yang dilakukan menggunakan data kecepatan *download*. Dari gambar tersebut terlihat bahwa *cluster* 0 memiliki jumlah anggota sebanyak 3 kelurahan, *cluster* 1 memiliki jumlah anggota sebanyak 13 kelurahan, *cluster* 2 dengan 5 kelurahan, dan *cluster* 3 dengan 35 kelurahan.



- Uji Fungsionalitas

Penelitian ini menggunakan metode *black box testing* untuk menguji aplikasi yang telah dibuat. Pengujian dengan metode *black box testing* digunakan untuk menguji suatu aplikasi tanpa harus memperhatikan secara detail dan hanya memeriksa keluaran berdasarkan masukan yang diberikan [37]. Uji fungsionalitas menggunakan metode *blackbox testing* pada penelitian ini dilakukan dengan menguji kesesuaian aksi yang dilakukan terhadap *dashboard* yang dibuat dimana hasil yang didapatkan pada lampiran penelitian ini. Hasil yang didapatkan pada uji fungsionalitas menunjukkan bahwa setiap aksi yang dilakukan pada *dashboard* yang dibuat telah sesuai dengan kebutuhan.

- Uji Validitas

Uji validitas pada penelitian ini menggunakan metode *content validity* untuk mengukur seberapa penting hal-hal yang ada pada *dashboard* untuk ditampilkan. Pengujian dilakukan dengan melakukan *interview* langsung terhadap 4 orang pakar untuk menilai suatu isi dari *dashboard* yang dibuat adalah sesuatu yang penting atau tidak. Tabel 21 memperlihatkan keempat orang pakar yang menjadi penguji pada uji validitas penelitian ini.

Tabel 27. Informasi pakar uji validitas

Pakar	Pakar 1	Pakar 2	Pakar 3	Pakar 4
Posisi	COO	Founder <i>Start Up</i>	Pengusaha Online	Pengusaha Online
Usia	44	44	35	31
Domisili	Kota Bekasi	Kota Bekasi	Kota Bekasi	Kota Bekasi

Tabel 21 menunjukkan hasil uji validitas yang dilakukan. Pengujian dilakukan dengan menguji 15 isi dari *dashboard* yang dibuat seperti yang terlihat pada tabel 4.19. Nilai 1 diberikan jika isi yang diuji dianggap penting dan nilai 0 diberikan jika isi yang diuji dianggap tidak penting. Nilai yang didapatkan dari setiap pakar

selanjutnya dijumlahkan dan dicari nilai rata-rata. Nilai rata-rata yang didapatkan dijumlahkan kembali dan dibagi 15 sesuai jumlah *item* yang diuji untuk mendapatkan nilai S-CVI. Nilai S-CVI yang diperoleh pada uji validitas yang dilakukan sebesar 0,88333. Nilai ini berada diatas nilai minimal yaitu 0,8 yang berarti *dashboard* yang dibuat telah lulus uji *content validity*.

Tabel 28. Hasil Uji Validitas dashboard

No	Item	Expert 1	Expert 2	Expert 3	Expert 4	I-CVI
1	Menu	1	1	1	1	1
2	Daftar isi	1	1	0	1	0,75
3	Grafik data <i>upload</i>	1	1	1	1	1
4	Tabel data <i>upload</i>	1	1	1	1	1
5	Filter data <i>upload</i>	1	1	1	1	1
6	Grafik data <i>download</i>	1	1	1	1	1
7	Tabel data <i>download</i>	1	1	1	1	1
8	Filter data <i>download</i>	1	1	1	1	1
9	Grafik data <i>upload all</i>	0	1	0	1	0,5
10	Tabel data <i>upload all</i>	1	1	1	1	1
11	Grafik data <i>download all</i>	0	1	0	1	0,5
12	Tabel data <i>download all</i>	1	1	1	1	1
13	Jumlah kelura- han berdasar- kan <i>cluster</i>	1	0	1	1	0,75

14	Identifikasi <i>cluster</i> setiap kelurahan	1	1	1	1	1
15	Grafik <i>cluster</i> untuk setiap kelurahan	1	1	1	0	0,75
S-CVI						0,88333

Model Business Intelligence (BI) yang dibuat menghasilkan informasi mengenai informasi deskriptif berupa kelurahan dengan kecepatan tertinggi dan kecepatan terendah baik untuk kecepatan upload maupun kecepatan download.

Kecepatan upload tertinggi sebesar 78Mbps saat menggunakan P2 di kelurahan Jatiwarna sedangkan kecepatan terendah sebesar 0,88Mbps saat menggunakan P3 di kelurahan Telukpucung. Kecepatan download tertinggi sebesar 100Mbps saat menggunakan P4 terdapat di kelurahan Pengasinan sedangkan kecepatan terendah saat menggunakan P2 sebesar 0,1Mbps saat menggunakan Telkom terdapat di kelurahan Jatimurni.

Pengelompokan wilayah kelurahan berdasarkan data upload terbagi menjadi 6 cluster dan pengelompokan berdasarkan data download dibagi menjadi 4 cluster berdasarkan nilai DBI.

Data upload dibagi kedalam 6 cluster dengan karakteristik cluster 0 memiliki nilai centroid P1, P2, dan P3 yang tinggi. Cluster 1 dengan keseluruhan provider bernilai centroid rendah. Cluster 2 memiliki nilai centroid tinggi pada P2. Centroid 3 memiliki nilai centroid yang tinggi pada P1 dan P3. Cluster 4 hanya P1 yang memiliki nilai centroid yang tinggi. Cluster 5 memiliki nilai centroid yang tinggi pada P2 dan P3.

Data download dibagi ke dalam 4 cluster dimana cluster 0 hanya memiliki P2 dengan nilai centroid yang tinggi, cluster 1 hanya memiliki P4 dengan nilai centroid yang tinggi, cluster 2 hanya memiliki P3 dengan nilai centroid yang tinggi, dan cluster 3 tidak memiliki nilai centroid yang tinggi.

Hasil clustering pada data upload menunjukkan cluster 0 memiliki jumlah anggota sebanyak 6 kelurahan, cluster 1 memiliki jumlah anggota sebanyak 20 kelurahan, cluster 2 dengan 12 kelurahan, cluster 3 dengan 8 kelurahan, cluster 4 dengan 7 kelurahan, dan cluster 5 dengan 3 kelurahan. Hasil clustering pada data download menunjukkan cluster 0 memiliki jumlah anggota sebanyak 3 kelurahan, cluster 1 memiliki jumlah anggota sebanyak 13 kelurahan, cluster 2 dengan 5 kelurahan, dan cluster 3 dengan 35 kelurahan.

Rancangan Business Intelligence (BI) yang dibuat mampu menghasilkan informasi-informasi yang terkait dengan kecepatan internet di setiap kelurahan kota Bekasi dan menampilkan hasil clustering yang dilakukan dan ditampilkan dalam prototype menggunakan Power BI. Prototype yang dibuat dinyatakan lolos uji fungsional dan dianggap sesuai oleh pakar dan lolos uji validitas dengan skor 0,8833.

DAFTAR PUSTAKA

- Anjaningrum W.D, Azizah N., Nanang Suryadi N. 2024. Spurring SMEs' performance through business intelligence, organizational and network learning, customer value anticipation, and innovation - Empirical evidence of the creative economy sector in East Java, Indonesia. *Heliyon* (10). <https://doi.org/10.1016/j.heliyon.2024.e27998>
- Aloraini, Fatimah, Amir Javed, Omer Rana, and Pete Burnap. 2022. Adversarial machine learning in IoT from an insider point of view. *Journal of Information Security and Applications* 70:103341. <https://doi.org/10.1016/j.jisa.2022.103341>.
- Atsani, M. R., Anjari, G. T., & Saraswati, N. M. 2019. Pengembangan Business Intelligence Di Rumah Sakit (Studi Kasus: RSUD Prof. Dr. Margono Soekarjo Purwokerto). *Telematika*, 12(2), 124–138. <https://doi.org/10.35671/telematika.v12i2.839>
- Harahap E. F., Fitriana R., Faturrahman C.L. 2021. Applying Data Mining of Association Rules as Decision Making In Coffee-Shop: a Case Study. *17th International Conference on Quality in Research (QIR): International Symposium on Electrical and Computer Engineering*. IEEE. p:66-70.
- Fitriana R, Eriyatno, Djatna T, Kusmuljono T.P. 2012. Peran Sistem Intelijensia Bisnis dalam Manajemen Pengelolaan Pelanggan dan Mutu untuk Agroindustri Susu Skala Usaha Menengah.

- Fitriana R. 2013. Rancang Bangun Sistem Intelijensia Bisnis untuk Agroindustri Susu Skala Menengah di Indonesia. Vol. Disertasi, Sekolah Pascasarjana Institut Pertanian Bogor.
- Fitriana R, Saragih J, and Hasyati BA. 2018. Perancangan Model Sistem Intelijensia Bisnis Untuk Menganalisis Pemasaran Produk Roti Di Pabrik Roti Menggunakan Metode Data Mining Dan Cube. *Jurnal Teknologi Industri Pertanian*. 28 (1): 113-126. 10.24961/j.tek.ind.pert.2018.28.1.113.
- Fitriana R, Saragih J, Luthfiana N. 2017. Model business intelligence system design of quality products by using data mining in R Bakery Company. *10th ISIEM. IOP Conf. Series: Materials Science and Engineering* 277. 012005:1-8.
- Fitriana R, Habyba A.N, Febriani E. 2022. Data Mining dan Aplikasinya. Banyumas: Wawasan Ilmu
- Hidayat, M. K., & Fitriana, R. 2022. Penerapan Sistem Intelijensia Bisnis Dan K-Means Clustering Untuk Memantau Produksi Tanaman Obat. *Jurnal Teknologi Industri Pertanian*, 32(2), 204–219.
- Iqbal, M. Z., Mustafa, G., Sarwar, N., Wajid, S. H., Nasir, J., & Siddque, S. (2020). A Review of Star Schema and Snowflakes Schema. *Communications in Computer and Information Science*, 1198, 129–140. https://doi.org/10.1007/978-981-15-5232-8_12
- Joshi, M., & Dubbawar, A. 2021. Review on Business Intelligence, Its Tools and Techniques, and Advantages and Disadvantages. *International Journal of Engineering Research & Technology*, 10(12), 386–391.
- Leiting, Ann K., Lien De Cuyper, and Christian Kauffmann. 2022. "The Internet of Things and the case of Bosch: Changing business models while staying true to yourself." *Technovation*, no. June 2020, 102497.
- Negash, S. 2004. Business Intelligence. *Communications of the Association for Information Systems*, 13. <https://doi.org/10.24961/j.tek.ind.pert.2018.28.1.113>

- Perry WE. 2006. *Effective Methods for Software Testing, Includes Complete Guidelines and Checklists*. Indiana: Wiley Publishing, Inc. Indianapolis
- Prasetyo, Wibowo H., Noor B. Naidu, Beti I. Sari, Rochman H. Mustofa, Naillysa Rahmawati, Gilang Pambudi A. Wijaya, and Obby T. Hidayat. 2021. "Survey data of internet skills, internet attitudes, computer self-efficacy, and digital citizenship among students in Indonesia." *Data in Brief* 39 (107569).
- Ramalingam S., Subramanian M., Reddy A.S., Abdullaev S., Dhahbi S. 2024. Exploring business intelligence applications in the healthcare industry: A comprehensive analysis. *Egyptian Informatics Journal* 25.
- Sargent, RG. 1999. *Validation and Verification of Simulation Model*. Proceeding of 1999 Winter Simulation Conference.
- Setiyani, L., Rostiani, Y., & Ratnasari, T. 2020. Analisis Kebutuhan Fungsional Sistem Informasi Persediaan Barang Perusahaan General Trading (Studi Kasus : PT. Amco Multitech). *Owner*, 4(1), 288. <https://doi.org/10.33395/owner.v4i1.205>
- Sharda, R., Delen, D., & Turban, E. 2018. *Business Intelligence, Analytics, And Data Science: A Managerial Perspective*. Pearson Education Limited.
- Sinaga, K. P., & Yang, M. S. 2020. Unsupervised K-means clustering algorithm. *IEEE Access*, 8, 80716–80727.
- Sugiarto, D., Mardianto, I., Najih, M., Adrian, D., & Pratama, D. A. (2021). Perancangan Dashboard Untuk Visualisasi Harga Dan Pasokan Beras Di Pasar Induk Beras Cipinang. *Jurnal Teknologi Industri Pertanian*, 31(1), 12–19.
- Thomason, Jane. 2022. "Metaverse, token economies, and non-communicable diseases." *Global Health Journal* 6 (3): 164-167.
- Turnip Y.E.H, Fitriana R. Clustering Kabupaten Berdasarkan Luas Hutan Menggunakan Metode K Means di Provinsi Jawa Tengah.

Jurnal Teknologi Industri Pertanian IPB. 33(1). 2022:1-9.

Vaisman, A., & Zimányi, E. (2014). Data warehouse systems: Design and implementation. In Data Warehouse Systems: Design and Implementation. Springer Berlin Heidelberg.

Vercellis, C. 2009. Business Intelligence: Data Mining and Optimization for Decision Making. In Business Intelligence: Data Mining and Optimization for Decision Making. <https://doi.org/10.1002/9780470753866>

Yusuf, B., Mahara, R., Ahmadian, H., Wahyuni, S., & AR, K. 2022. Analisis Clustering Penduduk Miskin Di Provinsi Aceh Menggunakan Algoritma K-Means Dan X-Means. Jurnal Nasional Komputasi Dan Teknologi Informasi, 5(1), 26–35.

Zhao J, Zang J. 2024. Machine Learning Based Business Intelligence Security And Privacy Analysis With Gaming Model In Training Complexity Application. Entertainment Computing. Volume 50,

PROFIL PENULIS



Rina Fitriana adalah Associate professor dan Dosen S1 Jurusan Teknik Industri, S2 Magister Teknik Industri dan S3 Doktor Teknik Industri FTI Universitas Trisakti. Rina memperoleh gelar Sarjana Teknik dari Jurusan Teknik Industri FTI Universitas Trisakti, gelar Magister Manajemen dari Sekolah Tinggi Manajemen PPM di Jakarta, gelar Doktornya Program Studi Teknologi Industri Pertanian Institut Pertanian Bogor dan gelar Insinyur dari Program Profesi Insinyur dari Unika Atma Jaya. Pada tahun 2017 sampai sekarang dipercaya menjadi Ketua Jurusan Teknik Industri FTI Universitas Trisakti. Beberapa mata kuliah yang diampu yaitu Pengendalian dan Penjaminan Mutu, Data Mining, Business Intelligence, Simulasi Sistem, Statistika Multivariat. Saat ini Rina aktif meneliti di bidang Rekayasa Kualitas, Sistem dan Simulasi Industri, Data Mining, Business Intelligence. Rina memiliki Sertifikasi Dosen, Certified International Business Intelligence Associate and Professional, Certified in Big Data Associate dan Insinyur Profesional Madya dan Asean Engineer. Rina juga telah dipercaya menjadi Asesor Beban Kerja Dosen. Rina telah membuat buku mengenai ajar mengenai Pengendalian Penjaminan Mutu dan Data Mining dan aplikasinya. Contoh kasus di industri manufaktur dan jasa. Rina dapat dihubungi pada email : rinaf@trisakti.ac.id.



Miwan Kurniawan Hidayat telah menyelesaikan pendidikan sarjana Teknik Informatika dari Universitas Pasundan Bandung pada tahun 2003, mendapatkan gelar Magister Ilmu Komputer dari Universitas Nusa Mandiri Jakarta pada tahun 2010, dan meraih gelar Magister Teknik Industri dari Universitas Trisakti Jakarta pada tahun 2023. Berprofesi sebagai praktisi Teknologi Informasi dan Dosen di Universitas

Bina Sarana Informatika dengan fokus pada bidang Rekayasa Informasi. Telah mendapatkan sertifikat profesi berupa Sertifikasi Dosen dan sertifikat kompetensi sebagai Asesor Kompetensi, *ICT Project Manager, Network Administrator, Programmer, Certificate in Business Intelligence Associate*. Miwan juga telah menerbitkan banyak makalah dalam jurnal ilmiah. Miwan dapat dihubungi melalui email: miwan@bsi.ac.id.



Yusri Eli Hotman Turnip adalah praktisi profesional di perusahaan global di bidang *healthcare*. Yusri menyelesaikan pendidikan Diploma 3 dari Jurusan Teknik Elektro Fakultas Teknik Universitas Gadjah Mada (UGM), menerima gelar Sarjana Teknik dari Jurusan Teknik Elektro Fakultas Teknik Universitas Pancasila, dan gelar Magister Teknik dari Jurusan Magister Teknik Industri FTI Trisakti. Memiliki

pengalaman bekerja di beberapa bidang seperti pertambangan, konstruksi bangunan, dan *healthcare* selama lebih dari 10 tahun. Yusri memiliki *Certified Business Intelligence Analyst (CBIA)* pada tahun 2022 untuk bidang *Business Intelligence*. Saat mengambil studi Magister, Yusri membuat publikasi di bidang *data mining* yang berjudul *District Clustering Based On Forest Area Using K-Means Method In Central Java Province* yang terbit pada jurnal Teknologi Industri Pertanian IPB.



Dr. Dedy Sugiarto, S.Si, M.M, M.Kom lahir pada tanggal 14 Oktober 1969 di Jakarta. Pendidikan dasar ditempuh di SDN Tebet Timur 19 Jakarta, pendidikan menengah di SMPN 115 Jakarta dan SMAN 8 Jakarta. Pendidikan sarjana ditempuh di Program Studi Statistika, FMIPA IPB Bogor, lulus tahun 1993. Pendidikan S2 di Program Studi

Magister Manajemen Universitas Trisakti, konsentrasi Manajemen Produksi/Operasi dengan beasiswa dari Universitas Trisakti, lulus tahun 1998. Pendidikan S3 di Program Studi Teknologi Industri Pertanian IPB Bogor dengan beasiswa Universitas Trisakti dan bantuan hibah penelitian mahasiswa program doktor dari Kementerian Pendidikan Nasional, lulus tahun 2012. Pendidikan S2 kembali ditempuh di Program Studi S2 PJJ Teknik Informatika Universitas AMIKOM Yogyakarta, lulus tahun 2023. Karir pertama kali sebagai dosen tetap di Jurusan Teknik Industri, Fakultas Teknologi Industri (FTI), Universitas Trisakti tahun 1994 dengan penempatan pada Laboratorium Statistika Industri. Tahun 2015 pindah *homebase* sebagai dosen tetap di Program Studi Sistem Informasi, FTI Universitas Trisakti dan membantu mengajar di beberapa program studi seperti Magister Teknik Industri FTI Trisakti, Magister Manajemen FEB Universitas Trisakti serta Program Studi Manajemen *Trisakti School of Management*. Saat ini menjabat sebagai sekretaris Jurusan Teknik Informatika FTI Universitas Trisakti sejak 2017, mengelola dua program studi di bawahnya yaitu Program Studi Informatika dan Program Studi Sistem Informasi. Sejak 2021 mendapatkan tugas dan tanggung jawab sebagai ketua bidang *data science, website* dan media sosial pada Gugus Tugas Implementasi dan Integrasi Teknologi Informasi Universitas Trisakti. Dedy dapat dihubungi melalui email dedy@trisakti.ac.id

SERTIFIKAT

Bukti Terbit Buku No: 567/Sert.Pnls-WI/VIII/2024

Diberikan Kepada :

Dr. Ir. Rina Fitriana, S.T., M.M., IPM.

Dr. Dedy Sugiarto, S.Si., M.M., M.Kom.

Miwan Kurniawan Hidayat, S.T., M.Kom., M.T.

Yusri Eli Hotman Turnip, S.T., M.T.

Judul Buku :

Perancangan Model Sistem Intelijensia Bisnis

ISBN :

978-623-132-304-0

**Yang Diterbitkan di Penerbit Wawasan Ilmu
Pada Tahun 2024**

Purwokerto, 4 September 2024
Direktur Penerbit Wawasan Ilmu

