

BAB III

PELAKSANAAN KERJA PROFESI

3.1 Bidang Kerja

Pada pelaksanaan kegiatan Kerja Profesi, Praktikan bekerja sebagai *Data Analyst* dalam divisi produk. Praktikan dipercayakan untuk melaksanakan tugas yang berkaitan dengan analisis data. Untuk mendukung kelancaran pengerjaan tugas, tools yang digunakan seperti *Google Colab* dengan bahasa pemrograman *Python*, *Metabase*, *BigQuery*, *Google Sheets*, dan *PowerPoint* untuk memproses, menganalisis, dan memvisualisasikan data. Selain itu, praktikan juga bertanggung jawab untuk menyajikan data dalam format yang mudah dipahami dan digunakan oleh tim *marketing*, serta memberikan wawasan yang mendukung keputusan strategis perusahaan dalam meningkatkan penjualan. Dalam melakukan analisis dan dokumentasi, praktikan melakukan beberapa pekerjaan sebagai berikut :

1. *User Behaviour Analysis* Menggunakan *BigQuery*
Praktikan melakukan analisis perilaku pengguna melalui *BigQuery*, yang memungkinkan untuk mengumpulkan, memproses, dan menganalisis data pengguna secara efisien. Data yang dianalisis mencakup pola perilaku pengguna pada produk Reksa dana untuk mengggali penyebab utama dari lonjakan *redeem* reksa dana di bulan April.
2. *Predictive Modelling* dengan Python untuk Prediksi Pengguna yang Berpotensi Membeli Produk
Praktikan menggunakan bahasa pemrograman Python dalam proyek *machine learning* untuk memprediksi pengguna yang berpotensi untuk membeli salah satu produk Bareksa yakni emas, berdasarkan pola perilaku mereka yang dianalisis sebelumnya. Hal ini bertujuan

untuk memberikan informasi yang lebih mendalam tentang peluang konversi dan memungkinkan pengembangan strategi pemasaran yang lebih efektif.

3. *Data Reporting* melalui Google Sheets dan PowerPoint Praktikan bertanggung jawab untuk menyajikan hasil analisis dalam format yang mudah dipahami oleh tim produk dan manajemen perusahaan. Hasil analisis disajikan melalui *Google Sheets* dan *PowerPoint* untuk memberikan gambaran yang jelas dan terstruktur tentang data dan prediksi yang telah dibuat.
4. Visualisasi Data untuk Memberikan Wawasan yang Mudah Dipahami Praktikan menggunakan *tools* seperti Metabase untuk memvisualisasikan data dan membuat grafik atau diagram yang memudahkan pemahaman dan pemantauan mengenai perilaku pengguna dan produk.

3.2 Pelaksanaan Kerja

Kegiatan berbentuk praktik dalam melaksanakan Kerja Profesi dimulai pada 10 Februari 2025 hingga 10 Mei 2025. Kerja Profesi ini diposisikan pada Divisi Produk sebagai Data Analisis untuk menganalisis dan memvisualisasikan data yang berkaitan dengan produk yang diperjual belikan di Bareksa. Selama masa kerja, praktikan sudah melakukan beberapa kegiatan dimulai dari pengenalan lingkungan kerja, *product knowledge*, berdiskusi dengan tim, melakukan pelatihan penggunaan *tools* pada perusahaan yaitu *BigQuery* dan *Metabase*, sampai pada tahap visualisasi data dan pelaporan melalui *google sheets*.

3.2.1 Data and Research Sync

Data and Research Sync merupakan pertemuan rutin setiap minggu yang dihadiri oleh seluruh anggota tim data dari divisi produk yang bersifat *online*, dipimpin oleh Gerry Fernando selaku *Senior*

Data and Research Manager yang bertujuan untuk menyelaraskan serta memperbarui perkembangan individu anggota tim dalam interval waktu satu minggu. Terdapat dua divisi dalam tim, yaitu *data science* dan *data analyst*, dan setiap individu memberikan presentasi lisan yang mencakup hasil analisis, kesimpulan, dan informasi lain yang telah diselesaikan dan dikerjakan. Setiap laporan yang dipresentasikan bertujuan untuk merangkum pekerjaan yang telah dilakukan secara spesifik dan juga untuk menjawab pertanyaan apakah hasil yang diperoleh akurat dan relevan dengan pekerjaan yang telah dilakukan. Hal ini memungkinkan anggota tim untuk berdiskusi, saling mengajarkan, dan memperbaiki masalah yang mungkin muncul selama pengerjaan. Dengan demikian, pertemuan ini bertujuan untuk memastikan apakah semua anggota memiliki pemahaman yang sama tentang tujuan utama dari tim data divisi

- produk sehingga informasi yang diperoleh dapat digunakan untuk meningkatkan kualitas dan efektivitas pengembangan produk.

Pada bagian kedua dari pertemuan, pekerjaan dibagi kepada setiap anggota tim secara individu. Pembagian tugas ini dilakukan sesuai dengan pekerjaan yang harus diselesaikan pada minggu mendatang dan dikategorikan sebagai hal yang sangat penting, memerlukan perhatian segera, serta berdasarkan ketersediaan dan keterampilan masing-masing anggota tim. Dengan cara ini, setiap anggota diberikan tanggung jawab yang sesuai dengan kemampuannya dan dijamin dapat diselesaikan tepat waktu. Tanggung jawab ini juga mencakup penetapan tenggat waktu yang dapat dicapai untuk memastikan bahwa target yang ditetapkan dapat dipatuhi dalam kerangka waktu yang telah ditentukan. Selain itu, pertemuan ini juga menjadi wadah untuk menyelesaikan masalah yang perlu ditangani dalam kegiatan yang sedang berlangsung dan memungkinkan kesepakatan bersama mengenai prinsip dan metode prosedur yang terkait dengan pengolahan data. Hal ini meningkatkan kesatuan di antara anggota tim dan lebih jauh lagi memastikan

bahwa hasil yang dihasilkan akurat dan memenuhi persyaratan pengembangan produk.

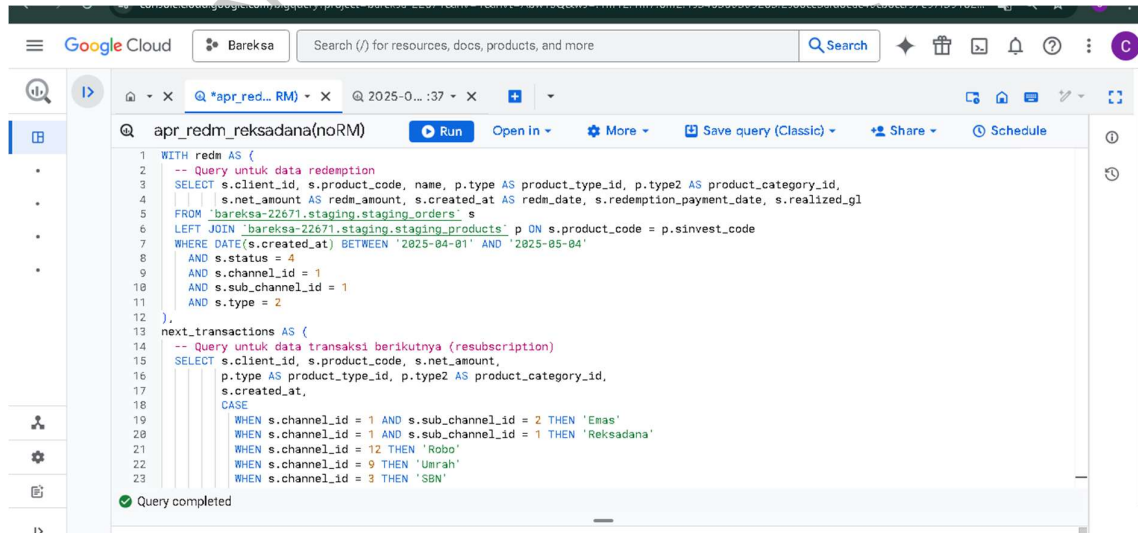
3.2.2 User Behaviour Analysis

Dalam menjalani kegiatan kerja profesi, Praktikan diberikan tugas untuk melakukan studi komprehensif tentang data pengguna yang terkait dengan salah satu produk Bareksa, khususnya produk reksa dana. Tujuan dari tugas ini adalah untuk memahami fenomena yang terjadi pada bulan April di mana terdapat lonjakan penjualan (*redemption*) pada bulan tersebut. *Redemption* adalah ketika seorang pengguna mencairkan dana dari produk investasi mereka, yang biasanya menunjukkan pergeseran strategi atau reaksi terhadap kondisi pasar. Praktikan juga diharuskan untuk menindaklanjuti pengguna yang kembali membeli produk (*resubscription*) setelah *redemption* untuk mencoba menentukan apakah ada pola atau tren tertentu yang mendorong keputusan pengguna untuk investasi kembali setelah menarik dana mereka. Karakteristik pengguna yang ingin dicari oleh Praktikan melalui analisis ini adalah pengguna yang kemungkinan besar melakukan *redemption* dan *subscription*, yang bersama dengan faktor-faktor lain akan dianalisis berkaitan dengan keputusan mereka. Selain itu, praktikan juga menghubungkan temuan ini untuk memberikan wawasan yang lebih lengkap, yang nantinya akan disajikan dalam laporan yang dapat digunakan oleh tim terkait, terutama tim marketing, untuk merancang strategi pemasaran dan peningkatan layanan produk yang lebih efektif.

1. Data Extraction and Processing

Untuk analisis ini, data dimulai dengan pengambilan dari *database* Bareksa menggunakan alat *query* seperti Google *BigQuery*. *Query* pertama bertujuan mendapatkan data *redemption* dengan menangkap bidang-bidang penting seperti ID pengguna,

kode produk, jumlah penjualan (*redemption*), nama produk, dan tanggal transaksi untuk pengguna yang menebus reksa dana mereka antara 01 April dan 4 Mei 2025. Seperti dalam kebanyakan skenario bisnis, hanya transaksi *redemption* yang valid yang dimasukkan sehingga *filter* tertentu pada status dan saluran diterapkan. Setelah itu, *query* lain dijalankan untuk mencakup data pembelian ulang(*resubs*) dan fokus beralih ke pengguna yang melakukan pembelian ulang ke produk setelah *redeem*. Berikut contoh query ekstraksi data pada **Gambar 3.1**, **Gambar 3.2** dan **Gambar 3.3**.



```

1 WITH redm AS {
2   -- Query untuk data redemption
3   SELECT s.client_id, s.product_code, name, p.type AS product_type_id, p.type2 AS product_category_id,
4         s.net_amount AS redm_amount, s.created_at AS redm_date, s.redemption_payment_date, s.realized_g1
5   FROM `bareksa-22671.staging.staging_orders` s
6   LEFT JOIN `bareksa-22671.staging.staging_products` p ON s.product_code = p.sinvest_code
7   WHERE DATE(s.created_at) BETWEEN '2025-04-01' AND '2025-05-04'
8         AND s.status = 4
9         AND s.channel_id = 1
10        AND s.sub_channel_id = 1
11        AND s.type = 2
12 },
13 next_transactions AS {
14   -- Query untuk data transaksi berikutnya (resubscription)
15   SELECT s.client_id, s.product_code, s.net_amount,
16         p.type AS product_type_id, p.type2 AS product_category_id,
17         s.created_at,
18         CASE
19           WHEN s.channel_id = 1 AND s.sub_channel_id = 2 THEN 'Emas'
20           WHEN s.channel_id = 1 AND s.sub_channel_id = 1 THEN 'Reksadana'
21           WHEN s.channel_id = 12 THEN 'Robo'
22           WHEN s.channel_id = 9 THEN 'Uarah'
23           WHEN s.channel_id = 3 THEN 'SBN'
24         END AS channel_name
25   FROM `bareksa-22671.staging.staging_orders` s
26   LEFT JOIN `bareksa-22671.staging.staging_products` p ON s.product_code = p.sinvest_code
27   WHERE DATE(s.created_at) BETWEEN '2025-04-01' AND '2025-05-04'
28         AND s.status = 4
29         AND s.channel_id = 1
30         AND s.sub_channel_id = 1
31         AND s.type = 2
32 }
33 SELECT * FROM redm, next_transactions

```

Gambar 3.1 Proses Ekstraksi Data di Bigquery (a)
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Bagian pertama *query* ini menggunakan *Common Table Expression* (CTE) yang disebut *redm*, data yang ditransformasi tanggal 4 April hingga 4 Mei diambil dari tahap *filtering* dari *staging_orders* yang digabungkan dengan *staging_products*. Proses penggabungan ini dilakukan dengan *cumulate status* = 4 yang berarti hanya transaksi yang berhasil yang akan dipilih. Selain itu, *filter channel_id* = 1 dan *sub_channel_id* = 1 digunakan untuk mengklasifikasikan pengguna pada produk reksa dana.

Penggabungan kolom *product_code* dari *staging_orders* dan *sinvest_code* dari *staging_products* berfungsi untuk mendeklarasikan kode produk yang akan didefinisikan oleh nama produk yang sesuai.

Bagian kedua dari *query* ini adalah CTE *next_transaction*, yang menangkap transaksi berikutnya setelah pengguna melakukan *redeem*, yang disebut sebagai pembelian ulang (*resubs*). Data yang diambil mencakup id pengguna, kode produk, jumlah transaksi, tipe produk, kategori produk, tanggal transaksi, dan produk yang dibeli kembali. Selain itu, klausa *CASE* dibuat untuk mendefinisikan jenis produk berdasarkan *channel_id* dan *sub_channel_id*. CTE ini hanya mencakup transaksi dengan status 4 dan tipe 1 (pembelian ulang) selama periode yang sama seperti sebelumnya.

```

25 END AS product_resubs
26 FROM `bareksa-22671.staging.staging_orders` s
27 LEFT JOIN `bareksa-22671.staging.staging_products` p ON s.product_code = p.sinvest_code
28 WHERE DATE(s.created_at) BETWEEN '2025-04-01' AND '2025-05-04'
29 AND s.status = 4
30 AND s.type = 1
31 AND s.channel_id IN (1, 3, 9, 12)
32 ),
33 user_identity AS (
34 -- Query untuk data identitas pengguna
35 SELECT DISTINCT client_id, uacc_id, sid, name, email, handphone, cc.created_at,
36 ROW_NUMBER() OVER (PARTITION BY client_id ORDER BY cc.created_at DESC) AS count_data
37 FROM `bareksa-22671.staging.staging_client_channels` cc
38 JOIN `bareksa-22671.staging.staging_clients` c USING(client_id)
39 WHERE cc.channel_id IN (1, 12, 9, 3)
40 AND c.account_source = 1
41 ),
42 df AS (
43 -- Gabungkan data redm, total redm, dan transaksi berikutnya
44 SELECT DISTINCT a.client_id, a.product_code, a.name AS product_name,
45 DATE(redm_date) AS redm_date, redm_amount, realized_gl,
46 total_redm_all_product, total_realized_gl,
47 STRING_AGG(DISTINCT product_resubs, ', ' ) AS product_resubs, SUM(net_amount) AS total_resubs

```

Query completed

Gambar 3.2 Proses Ekstraksi Data di Bigquery (b)
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

CTE ketiga, *user_identity*, berfungsi untuk mengambil informasi pengguna yang relevan dengan transaksi. Data yang diambil mencakup *client_id*, *uacc_id*, *sid*, nama pengguna, email, nomor telepon, dan tanggal pembuatan akun (terkait dengan saluran

spesifik). *ROW_NUMBER()* digunakan untuk memberikan nomor urut pada setiap *client_id* berdasarkan urutan tanggal pembuatan akun secara menurun, yang menjamin bahwa hanya satu entri per pengguna yang akan dipilih.

```

45  DATE(redm_date) AS redm_date, redm_amount, realized_gl,
46  total_redm_all_product, total_realized_gl,
47  STRING_AGG(DISTINCT product_resubs, ', ' ) AS product_resubs, SUM(net_amount) AS total_resubs
48  FROM redm a
49  LEFT JOIN (
50    SELECT client_id, SUM(redm_amount) AS total_redm_all_product, SUM(realized_gl) AS total_realized_gl
51    FROM redm
52    GROUP BY client_id
53  ) ab USING(client_id)
54  LEFT JOIN next_transactions b ON a.client_id = b.client_id AND b.created_at > redm_date
55  GROUP BY a.client_id, a.product_code, a.name, redm_date, redm_amount, realized_gl, total_redm_all_product, total_realized_gl
56  )
57  SELECT df.client_id, uacc_id, ab.email, handphone, rm_name, product_code, product_name, redm_date, redm_amount, realized_gl,
58  total_redm_all_product, total_realized_gl, product_resubs, total_resubs,
59  ROW_NUMBER() OVER (PARTITION BY df.client_id ORDER BY redm_date) AS num_data,
60  CASE
61    WHEN total_resubs > total_redm_all_product THEN 'Resubs > Redm'
62    ELSE 'Resubs <= Redm'
63  END AS resub_vs_redm
64  FROM df
65  LEFT JOIN user_identity ab ON df.client_id = ab.client_id AND ab.count_data = 1
66  LEFT JOIN 'bareksa-22671-staging.nasabah_rm' rm USING(sid)
67  ORDER BY df.client_id, redm_date;

```

Query completed

Gambar 3.3 Output Ekstraksi Data di Bigquery (a)
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

CTE *df* adalah langkah integrasi yang menghubungkan informasi dari transaksi penukaran (*redm*), total penukaran di setiap produk (*total_redm_all_product*), bersama dengan data transaksi pembelian ulang berikutnya. Di sini, data penukaran digabungkan dengan total penukaran dan informasi transaksi untuk transaksi selanjutnya. Data produk yang terlibat dalam pembelian ulang digabungkan menjadi satu kolom melalui penggunaan *STRING_AGG* untuk menyederhanakan analisis. Selain itu, jumlah total pembelian ulang (*total_resubs*) per pengguna dihitung.

Setelah menggabungkan data di CTE *df*, bagian akhir dari *query* ini mengambil informasi lebih rinci seperti id pengguna, identifikasi pengguna, nama RM, produk, tanggal penukaran, jumlah penukaran, dan total pembelian ulang. Hasil keluaran dari *query* ini

juga menambahkan kolom baru *resub_vs_redm* yang menghitung perbandingan antara pembelian ulang yang terintegrasi dengan jumlah penukaran yang terintegrasi per pengguna, mencatat apakah transaksi pembelian ulang melebihi penukaran atau tidak. Untuk setiap pengguna, *ROW_NUMBER()* diterapkan untuk mengurutkan secara berurutan berdasarkan tanggal pembelian terbaik. Hasilnya diurutkan dengan *client_id* dan tanggal penukaran. Berikut contoh hasil dari *query* yang dilakukan terdapat pada **Gambar 3.4**.



The image shows a screenshot of a BigQuery query result. The table has several columns, including 'client_id', 'transaction_date', and 'resub_vs_redm'. The data is sorted by 'client_id' and 'transaction_date'. The 'resub_vs_redm' column contains numerical values representing the ratio of resubscriptions to redemptions for each user.

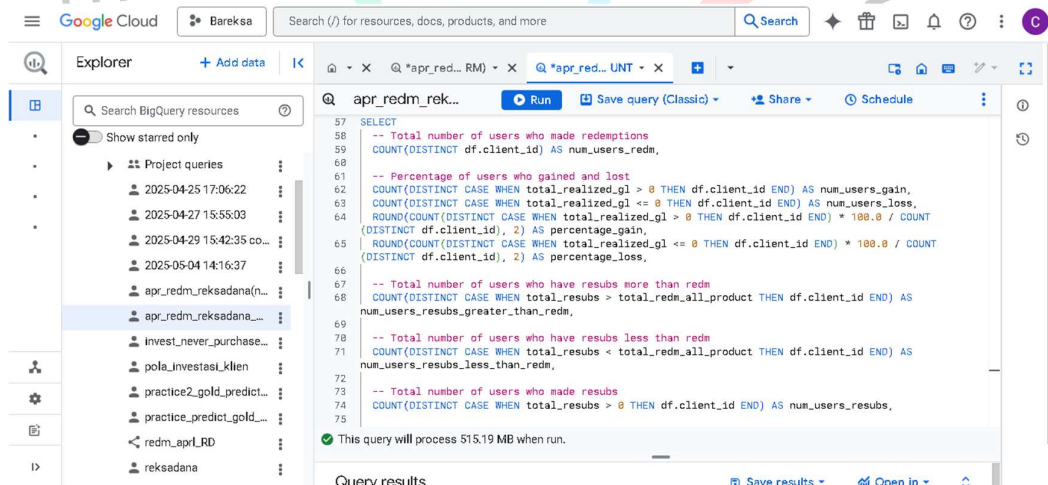
Gambar 3.4 Output Ekstraksi Data di Bigquery (b)
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Output yang dihasilkan oleh *query* ini mencakup detail relevan mengenai pengguna yang melakukan pembelian investasi baru pada produk-produk seperti reksa dana, emas, dan produk keuangan lainnya setelah melakukan *redeem*. Selain itu, *query* juga dijalankan untuk mengidentifikasi pengguna dengan tujuan menangkap informasi yang diperlukan, memastikan terjadinya keterkaitan yang tepat antara *redeem events* dan *resubscription events*, serta melengkapi pengguna dengan informasi seperti nama, email, no. telepon, dan status akunya. Dataset yang diekstrak kemudian digabung menjadi satupadu berdasarkan atribut yang relevan untuk menggambarkan secara detail aktivitas *redeem* and *resubscription* dari pengguna. Hal ini memungkinkan analisis tren

yang lebih mendalam dan memberikan wawasan lebih jauh mengenai perilaku transaksi dan pola investasi pengguna

2. Trend Analysis

Setelah mengumpulkan data relevan dari berbagai CTE seperti *redm* untuk transaksi penjualan (*redemption*), *next_transactions* untuk transaksi pembelian ulang, dan *user_identity* untuk identitas pengguna, data tersebut diintegrasikan ke dalam *df* CTE untuk pemeriksaan lebih lanjut terkait interaksi antara *redemption* dan *resubscription* secara terperinci. Pada tahap *query* ini, beberapa metrik tentang perilaku pengguna dibuat. Salah satu metrik tersebut berkaitan dengan penghitungan pengguna yang melakukan *redemption* dan menghitung persentase pengguna yang memperoleh atau kehilangan berdasarkan nilai *realize_gl*. Berikut contoh *query trend analysis* yang terdapat pada **Gambar 3.5** dan **Gambar 3.6**.

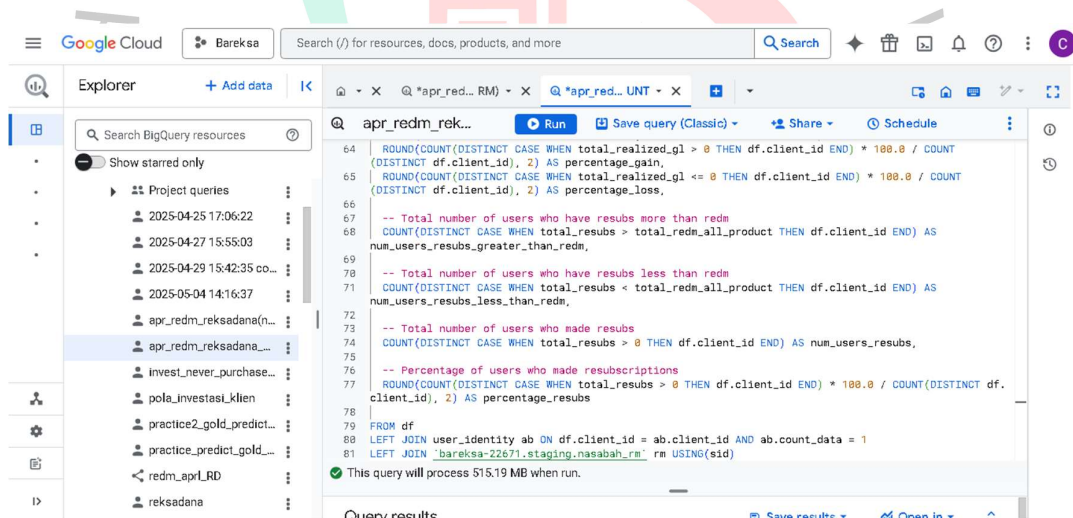


Gambar 3.5 Trend Analysis pada BigQuery (a)

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Metrik-metrik penting dihimpun dan dihitung untuk memperoleh wawasan terkait perilaku pengguna, khususnya dalam

kaitannya dengan transaksi *redemption* dan *resubscription*. Pertama-tama, *query* ini menghitung jumlah pengguna yang melakukan *redemption*. Untuk keperluan ini, fungsi `COUNT(DISTINCT df_user_level.client_id)` digunakan untuk mendapatkan jumlah unik pengguna yang melakukan transaksi *redemption* sehingga tidak ada duplikasi pengguna pada kalkulasi. Selanjutnya, dilakukan analisis terhadap pengguna yang mengalami keuntungan atau kerugian berdasarkan nilai `realized_gl`. Pengguna dengan nilai `realized_gl` positif dianggap memperoleh keuntungan, sedangkan mereka yang memiliki nilai `realized_gl` negatif atau nol dianggap mengalami kerugian. Untuk memberikan gambaran yang lebih menyeluruh, proporsi pengguna yang mengalami keuntungan atau kerugian juga dihitung dengan membagi jumlah pengguna dalam masing-masing kategori terhadap total pengguna *redemption*, kemudian dikalikan dengan 100.



Gambar 3.6 Trend Analysis pada BigQuery (b)
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Query ini juga menganalisis hubungan antara transaksi *resubscription* dan *redemption* di antara para pengguna. Analisis ini dilakukan untuk mengetahui sejauh mana transaksi *resubscription* pengguna melebihi, kurang dari, atau sama dengan jumlah

redemption yang mereka lakukan. Ini dicapai melalui pernyataan *CASE* di mana *total_resubs* dibandingkan dengan *total_redm_all_product*, dan pengguna yang memenuhi kondisi tersebut dihitung. Sebagai contoh, *num_users_resubs_greater_than_redm* menghitung pengguna yang melakukan *resubscription* lebih besar dari *redemption* mereka, *num_users_resubs_less_than_redm* menghitung pengguna yang melakukan *resubscription* lebih kecil dari jumlah *redemption*, dan *num_users_resubs_equal_redm* menghitung pengguna yang jumlah *resubscription*-nya sama dengan *redemption*. Dengan cara ini, *query* tidak hanya menyediakan angka terkait transaksi, tetapi juga menawarkan pemahaman yang lebih mendalam tentang hubungan pengguna dengan produk, yang dapat dimanfaatkan dalam strategi pemasaran dan manajemen produk di masa mendatang.

Rasio antara pengguna yang melakukan *resubscription* dan total pengguna yang melakukan *redemption* juga dihitung untuk memperkirakan proporsi pengguna *redemption* yang melanjutkan ke transaksi *resubscription*, terlepas dari apakah jumlahnya lebih besar, lebih kecil, atau sama. Tidak seperti kelompok sebelumnya, pengguna ini hanya dihitung jika mereka memiliki transaksi *resubscription* dalam bentuk apa pun. Nilai *num_users_resubs* mencerminkan metrik tersebut, dengan perhitungan persentase dilakukan dengan pendekatan yang sama. Berikut hasil *query trend analysis* pada **Gambar 3.7**

Query results

Save results

Open in

Job information

Results

Chart

JSON

Execution details

Execution graph

Rgm_users_redm	num_users_gain	num_users_loss	percentage_gain	percentage_loss	num_users_resubs_g	num_users_resubs_l	
1	14812	13169	1643	88.91	11.09	2227	3051

Gambar 3.7 Output Trend Analysis pada BigQuery (a)

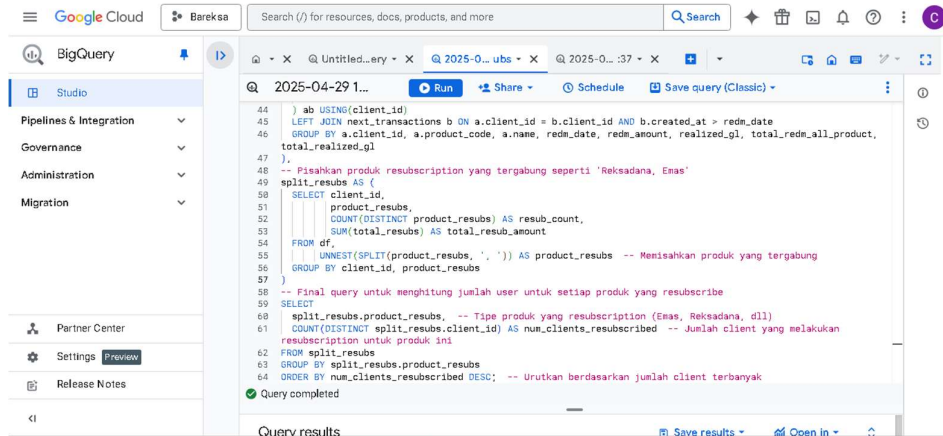
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Dengan menggabungkan semua metrik yang dijelaskan, kita memperoleh pemahaman yang lebih mendalam tentang interaksi pengguna terhadap aktivitas *redemption* dan *resubscription*, serta hubungan antara kedua aktivitas tersebut. Ini memberikan wawasan tentang perilaku pengguna dalam bertransaksi, khususnya apakah mereka cenderung melakukan *resubscription* setelah *redemption*, serta bagaimana nilai yang mereka peroleh dari *redemption* memengaruhi keputusan transaksi berikutnya.

Untuk mengidentifikasi produk yang paling banyak *di-resubscription*, Praktikan mengakumulasi data pengguna yang melakukan *resubscription*. Isu utama yang menjadi fokus dalam *query* ini adalah pemisahan produk *resubscription* yang sebelumnya tergabung, seperti “Reksadana, Emas”, agar dapat diperlakukan sebagai entri terpisah dalam menghitung partisipasi pengguna.

Pada bagian pertama, sebuah *Common Table Expression* (CTE) dengan alias *split_resubs* memproses data produk terkait dengan *resubscription*. Data produk, atau data *resubscription*, disimpan dalam tabel bernama *df* yang menyimpan catatan transaksi *resubscription* pengguna. Produk yang sebelumnya ditulis dalam satu kolom *product_resubs* sebagai “Reksadana, Emas” dan produk lainnya diubah menjadi entri terpisah menggunakan fungsi *SPLIT()*. Hasil dari pemisahan tersebut kemudian diproses lebih lanjut menggunakan fungsi *UNNEST()* yang menghasilkan setiap produk yang tergabung sebelumnya menjadi baris yang terpisah, sehingga setiap produk yang telah digabung kini ditampilkan sebagai entri yang terpisah berdasarkan *client_id*. Agregasi juga dilakukan untuk menentukan nilai untuk setiap produk *resubscription* yang terpisah. Jumlah produk yang terhitung disebut sebagai *resub_count* dan total transaksi *resubscription* untuk produk yang dihitung tersebut disebut sebagai *total_resub_amount*, keduanya dikelompokkan berdasarkan *client_id* dan produk *resubscription*, serta dihitung secara terpisah

untuk setiap produk dan menghitung jumlahnya. Berikut *query* pada **Gambar 3.8**.



Gambar 3.8 Query Trend Analysis pada BigQuery
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Bagian kedua dari *query* ini melakukan agregasi tambahan untuk menghitung jumlah pengguna unik (`num_clients_resubscribed`) yang melakukan *resubscription* untuk setiap jenis produk yang telah dipisahkan. Hasilnya adalah daftar produk yang telah *di-resubscribe*, di mana setiap produk memiliki jumlah pengguna yang melakukan *resubscription* untuk produk tersebut. *Output* ini kemudian diurutkan berdasarkan jumlah total pengguna yang melakukan *resubscription*, dengan produk yang paling sering *di-resubscribe* ditampilkan terlebih dahulu. Berikut contoh output produk yang paling sering *di-resubscribe* pada **Gambar 3.9**.

Query results

Save results Open in

Job information Results Chart JSON Execution details Execution graph

Row	product_resubs	num_clients_resubsc
1	Reksadana	1000
2	Emas	1000
3	SBN	1000
4	Robo	1000
5	Umrah	1000

Gambar 3.9 Output Trend Analysis pada BigQuery (b)

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Dengan analisis produk tersebut, dapat menentukan produk mana yang paling sering *di-resubscribe* oleh pengguna, serta partisipasi pengguna untuk setiap produk tersebut. Query ini membantu dalam mempelajari pola aktivitas *resubscription* pengguna sehubungan dengan produk-produk yang ada dan membantu dalam strategi pemasaran dan pengembangan produk.

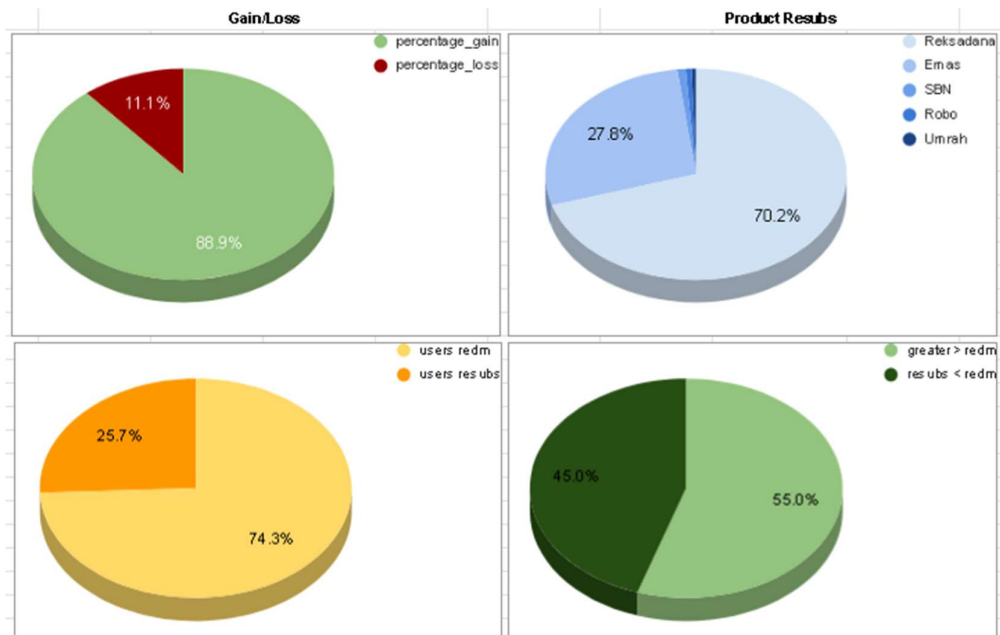
3. Reporting

Agar hasil dari *query* tersebut dapat tersajikan dengan baik, Praktikan menyajikan hasil *query* ke *google sheet* dan divisualisasikan. Seperti pada **Gambar 3.10** dan **Gambar 3.11**.

Product	Client ID	Subscription Date	Product Name	Subscription Type	Subscription Status	Subscription Duration	Subscription Frequency	Subscription Amount	Subscription Location	Subscription Channel	Subscription Source	Subscription Target
Reksadana	1000	2020-01-01	Reksadana	Subscription	Active	12 Months	Monthly	1000	Indonesia	Online	Google	Google
Emas	1000	2020-01-01	Emas	Subscription	Active	12 Months	Monthly	1000	Indonesia	Online	Google	Google
SBN	1000	2020-01-01	SBN	Subscription	Active	12 Months	Monthly	1000	Indonesia	Online	Google	Google
Robo	1000	2020-01-01	Robo	Subscription	Active	12 Months	Monthly	1000	Indonesia	Online	Google	Google
Umrah	1000	2020-01-01	Umrah	Subscription	Active	12 Months	Monthly	1000	Indonesia	Online	Google	Google

Gambar 3.10 Reporting pada Google Sheets (a)

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan



Gambar 3.11 Visualisasi pada Google Sheets

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Secara keseluruhan, informasi ini menunjukkan bahwa ada lebih banyak pengguna yang melakukan *redemption* dibandingkan dengan yang melakukan *resubscription*. Dalam pandangan yang lebih teragregasi, dapat dilihat bahwa beberapa produk investasi, termasuk reksadana, memiliki tingkat *resubscription* yang secara signifikan lebih rendah atau, pada beberapa kasus, saldo antara *redemption* dan *resubscription* yang saling tumpang tindih. Hal ini juga menunjukkan bahwa persentase klien yang melakukan *resubscription* masih rendah, dengan jumlah pengguna yang mengalami kerugian (*loss*) sebesar 11,09% dan pengguna yang mengalami keuntungan (*gain*) sebesar 88,91%. Pada dasarnya, mayoritas pengguna yang terlibat lebih memilih untuk menarik investasi mereka (*redemption*) saat mengalami keuntungan (*gain*) daripada meningkatkan investasi secara eksponensial dan melakukan *resubscription* pada produk yang sama, yang

menunjukkan ketidakpastian atau kurangnya kepercayaan terhadap kinerja produk-produk tersebut.

Selain itu, distribusi pengguna berdasarkan jenis produk menunjukkan keberagaman yang signifikan, dengan reksadana mendominasi jumlah pengguna yang melakukan *resubscription*. Hal ini tercermin dari 70,12% pengguna yang berpartisipasi dalam *resubscription* dengan produk reksadana, diikuti dengan produk emas, surat berharga negara (SBN), robo, dan umrah. Pengguna yang melakukan *resubscription* menunjukkan permintaan yang kuat terhadap produk investasi yang lebih aman dan terdiversifikasi. Produk-produk yang lebih stabil dan kurang berisiko, seperti emas dan SBN, cenderung menarik investor yang ingin menghindari fluktuasi pasar yang tajam. Di sisi lain, produk robo dan umrah, meskipun jumlahnya lebih sedikit, juga menunjukkan adanya minat dari klien yang mungkin mencari diversifikasi di saluran investasi alternatif lainnya.

3.2.3 Predictive Modelling

Dalam melaksanakan tugas ini, Praktikan diberikan peran untuk melakukan evaluasi menyeluruh terhadap data pengguna yang berkaitan dengan produk emas Bareksa, dengan kondisi pengguna ini tidak memiliki riwayat transaksi sebelumnya dengan produk emas. Tujuan dari analisis ini adalah untuk mengklasifikasikan pengguna yang berpotensi memiliki keterlibatan tinggi dengan produk emas di masa depan, meskipun mereka belum pernah melakukan transaksi emas. Terkait dengan investasi pada produk emas, Untuk mengetahui kemungkinan minat berinvestasi pengguna, Praktikan menganalisis transaksi produk diluar emas, pola pembelian pengguna, frekuensi dan keteraturan pengguna dalam bertransaksi, serta status portofolio pengguna. Pada fase ini, Praktikan mencoba untuk menemukan aspek-aspek yang dapat mendorong pengguna untuk berinvestasi, tanpa memandang apakah mereka telah membeli produk lain atau tidak memiliki riwayat pembelian emas, dan

menganalisis faktor-faktor tersebut untuk menarik kesimpulan. Analisis ini akan membantu departemen pemasaran untuk mendefinisikan ulang strategi mereka yang ditujukan kepada klien yang berpotensi berinvestasi dalam emas serta meningkatkan efektivitas kampanye pemasaran dan layanan produk emas yang ditawarkan oleh Bareksa.

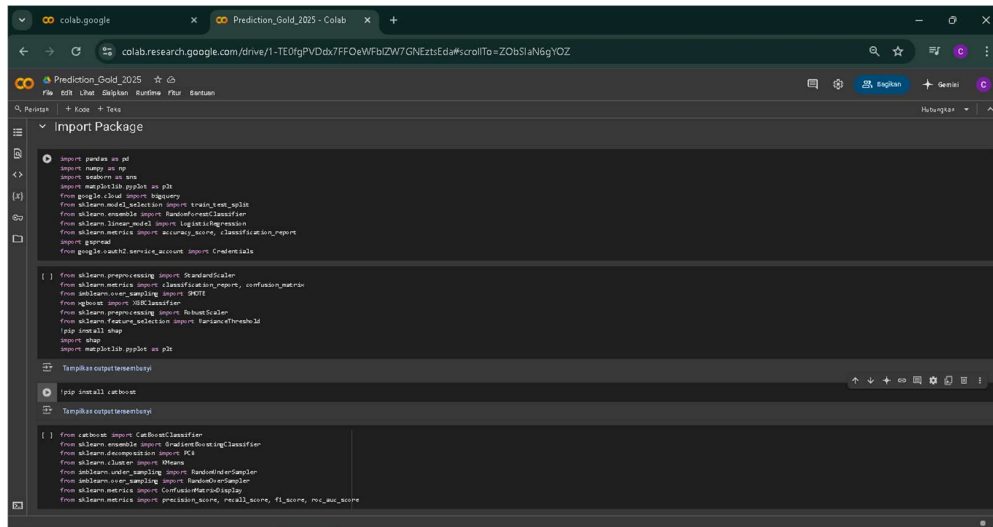
1. Data Preparation and Importing Required Package

Proses ini dimulai dengan mengimpor pustaka dan alat yang diperlukan untuk analisis data dan pemodelan prediktif. Seperti biasa, pustaka `pandas` dan `numpy` yang penting sangat krusial untuk pemrosesan dan manipulasi data karena mereka memfasilitasi penanganan kumpulan data besar. Untuk visualisasi data, digunakan `matplotlib` dan `seaborn` sebagai alat pemrograman untuk menciptakan representasi visual yang rinci dari data yang tersedia.

Sejumlah pustaka pembelajaran mesin lainnya diterapkan untuk membangun dan mengevaluasi model prediktif. Salah satu sumber daya penting adalah pustaka `sklearn`, yang memberikan akses ke banyak algoritma pengembangan model seperti `RandomForestClassifier`, `LogisticRegression`, dan `XGBClassifier` yang digunakan dalam berbagai tugas prediktif. Untuk mengevaluasi model pada data yang tidak tersentuh, kumpulan data dibagi menjadi set pelatihan dan pengujian menggunakan fungsi `train_test_split` yang ditemukan di `sklearn.model_selection`.

Untuk menghilangkan masalah ketidakseimbangan kelas yang dapat menghambat efektivitas model, pustaka `imblearn` digunakan. Menggabungkan langkah-langkah ini membantu mencapai generalisasi yang lebih baik untuk semua kelas dalam data yang disediakan. Pustaka-pustaka ini meningkatkan keterbacaan, otomatisasi, dan kepercayaan selama tahap pemodelan yang mengarah pada prediksi yang lebih maju, misalnya, mendeteksi pengguna dengan kecenderungan meningkat untuk berinvestasi

dalam produk emas di tahun-tahun mendatang. Berikut seperti pada **Gambar 3.12**.



```
import pandas as pd
import numpy as np
import sklearn as sk
import matplotlib.pyplot as plt
from google.cloud import bigquery
from sklearn.model_selection import train_test_split
from sklearn.ensemble import RandomForestClassifier
from sklearn.linear_model import LogisticRegression
from sklearn.metrics import accuracy_score, classification_report
import sys

from sklearn.preprocessing import StandardScaler
from sklearn.metrics import classification_report, confusion_matrix
from sklearn.cross_validation import train_test_split
from sklearn.metrics import roc_auc_score
from sklearn.preprocessing import RobustScaler
from sklearn.feature_selection import RFE
import shap
import time
import matplotlib.pyplot as plt

!pip install catboost

from catboost import CatBoostClassifier
from sklearn.ensemble import GradientBoostingClassifier
from sklearn.decomposition import PCA
from sklearn.cluster import KMeans
from sklearn.cross_validation import train_test_split
from sklearn.metrics import roc_auc_score
from sklearn.metrics import ConfusionMatrixDisplay
from sklearn.metrics import precision_score, recall_score, f1_score, matthews_corrcoef
```

Gambar 3.12 Importing Required Packages

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Setelah melakukan *import package*, Praktikan menggunakan data yang berasal dari BigQuery untuk melatih model *machine learning*. *Query* ini bertujuan untuk melakukan analisis mendalam terhadap data pengguna untuk memperkirakan kemungkinan mereka membeli produk emas, tanpa adanya transaksi emas sebelumnya. Langkah pertama dalam proses ini adalah mengidentifikasi pengguna yang tidak pernah membeli emas. Hal ini dilakukan dengan menggunakan *Common Table Expression* (CTE) `tbl_user_non_emas`, yang mencari pengguna dengan status dan kategori tertentu dalam tabel transaksi sebagai pengguna yang tidak melakukan pembelian emas, disaring berdasarkan status, kategori produk, *channel* transaksi, dan status pengguna. Pada CTE ini, Praktikan menarik pengguna yang melakukan transaksi pada produk reksa dana dengan rentang waktu pembuatan akun 2022 sampai dengan 31 Maret 2025. Setelah pengguna yang tidak membeli emas diidentifikasi, langkah berikutnya adalah menentukan tanggal pertama kali mereka membeli emas. Untuk tujuan ini, CTE

tbl_first_emas dibuat untuk mengekstrak data mengenai tanggal pembelian emas pertama dari tabel transaksi yang relevan. Hal ini membantu menetapkan dasar bagi pengguna yang mungkin mempertimbangkan untuk membeli emas di masa depan, meskipun mereka belum pernah melakukan transaksi emas sebelumnya. Proses ini memastikan bahwa setiap pengguna yang sudah bertransaksi emas tercatat dan dipertimbangkan dalam analisis.

Selanjutnya, CTE *tbl_limitation* digunakan untuk menggabungkan informasi dari *tbl_user_non_emas* dan *tbl_first_emas* untuk menetapkan tanggal batasan atau *date_limitation* bagi setiap pengguna yang terpilih. Bagi pengguna yang belum pernah membeli emas, tanggal pembatasnya ditetapkan pada 31 Maret 2025. Namun, bagi pengguna yang sudah membeli emas, tanggal pembatasnya disesuaikan dengan tanggal pembelian emas pertama. Pembatasan ini memungkinkan analisis lebih terfokus yang berpusat pada transaksi yang terjadi sebelum pembelian emas pertama atau dalam rentang waktu yang relevan. CTE *tbl_order* digunakan untuk menangkap transaksi lain di luar produk emas, termasuk *subscription* dan *redemption*, untuk menganalisis aktivitas pengguna terkait produk non-emas. Data dari tabel ini dibatasi pada transaksi dengan status selesai (*status=4*) dan tipe transaksi berupa *subscription* (*tipe = 1*) atau *redemption* (*tipe = 2*). Pola informasi ini membantu menganalisis transaksi pengguna terkait produk keuangan lain yang bukan emas, yang dapat membantu mengetahui kemungkinan minat mereka terhadap produk emas di masa depan.

Untuk menganalisis lebih lanjut perilaku transaksi pengguna, metrik transaksi didefinisikan berdasarkan pengguna dan bulan dalam CTE *tbl_order_month*. Di CTE ini, total nilai transaksi bulanan (*net amount*) dan jumlah hari transaksi per bulan dihitung. CTE ini memberikan gambaran pola pembelian yang terstruktur dan rutin, yang sangat berguna untuk segmentasi pengguna berbasis transaksi

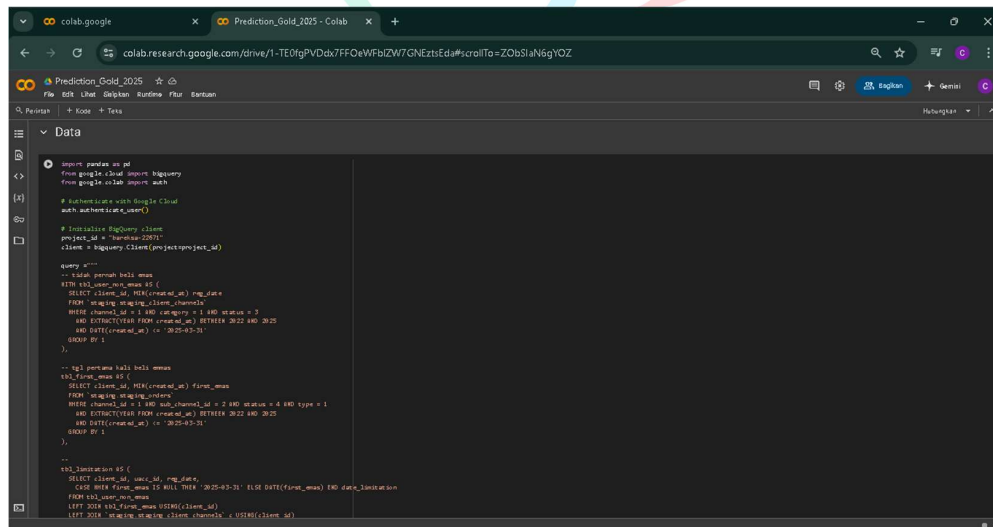
untuk menargetkan pengguna yang berpotensi tertarik pada produk investasi lainnya, termasuk emas, berdasarkan pola transaksi mereka. CTE `tbl_order_interval` dirancang untuk mempelajari waktu, yang diukur dengan fungsi LAG, antara dua transaksi berturut-turut. Dengan mengukur waktu antara transaksi yang berurutan, analisis ini membantu mengidentifikasi pengguna dengan pola transaksi yang berulang. Pengguna yang memiliki perilaku transaksi yang teratur lebih mungkin untuk melanjutkan investasi pada produk lain, seperti emas, karena mereka memiliki kecenderungan investasi yang konsisten.

Selanjutnya, CTE `tbl_rutin` menghitung rata-rata interval antara transaksi untuk menganalisis seberapa sering pengguna melakukan transaksi dalam waktu tiga bulan atau kurang. Pengguna yang bertransaksi secara rutin (dan lebih spesifik lagi, yang bertransaksi lebih sering) mendapatkan status transaksi yang telah ditentukan, yang membantu mengukur sejauh mana partisipasi mereka dalam investasi yang stabil. Pengguna dengan pola transaksi yang rutin biasanya lebih cenderung untuk berinvestasi lebih lanjut, termasuk dalam produk emas. CTE `tbl_interval_redm` memberikan pemahaman mengenai perilaku pengguna dengan menganalisis interval antara redemption dan penarikan dana dari produk investasi mereka. Pengguna yang sering melakukan redemption mungkin juga tertarik pada produk emas untuk diversifikasi portofolio mereka.

Informasi mengenai transaksi *subscription* dan *redemption* terakhir dari setiap pengguna disimpan dalam CTE `tbl_last_order`. Data ini sangat penting untuk memantau aktivitas terakhir pengguna dalam hal investasi dan memberikan gambaran terbaru tentang keputusan mereka untuk berinvestasi kembali atau menarik dana. Hal ini membantu memahami siklus investasi pengguna dan memprediksi kecenderungan mereka untuk membeli emas di masa depan. Selain itu, CTE `tbl_porto` mengambil data portofolio

pengguna, yang mencakup total aset yang dikelola (AUM). Dengan mengetahui nilai portofolio dan kemampuan finansial pengguna, kita dapat menilai potensi mereka untuk berinvestasi lebih lanjut dalam produk yang lebih stabil seperti emas. Ini memberikan wawasan penting tentang kemampuan finansial mereka untuk membeli emas.

Terakhir, semua data yang telah dihimpun dari berbagai CTE ini digabungkan dalam CTE utama df, yang menghitung berbagai fitur penting seperti total subscription (subs), redemption (redm), frekuensi transaksi, serta rata-rata interval transaksi dan redemption. Pada tahap ini, label tambahan (label) ditambahkan untuk menandai apakah pengguna layak untuk dianalisis lebih lanjut, dan dataset ini diurutkan berdasarkan total subscription. Dengan menggunakan data ini, analisis mendalam dilakukan untuk mengidentifikasi pengguna yang mungkin tertarik membeli emas di masa depan, bahkan tanpa riwayat pembelian emas sebelumnya. Berikut contoh *code* python yang dijalankan dalam *notebook* pada **Gambar 3.13** dan contoh output dataframe pada **Gambar 3.14**.



```

import pandas as pd
from google.cloud import bigquery
from google.colab import auth

# Authenticate with Google Cloud
auth.authenticate_user()

# Initialize BigQuery client
project_id = 'hematop-2025'
client = bigquery.Client(project=project_id)

query = """
-- table pertama hasil emas
WITH tbl_user_mas AS (
  SELECT client_id, MIN(created_at) mg_date
  FROM `bigquery-public-data.crypto_transactions`
  WHERE channel_id = 1 AND category = 1 AND status = 3
  AND EXTRACT(YEAR FROM created_at) BETWEEN 2022 AND 2025
  AND DATE(created_at) <= '2025-03-31'
  GROUP BY 1
),
-- table kedua hasil emas
tbl_firm_mas AS (
  SELECT client_id, MIN(created_at) firm_date
  FROM `bigquery-public-data.crypto_transactions`
  WHERE channel_id = 1 AND sub_channel_id = 2 AND status = 4 AND type = 1
  AND EXTRACT(YEAR FROM created_at) BETWEEN 2022 AND 2025
  AND DATE(created_at) <= '2025-03-31'
  GROUP BY 1
),
-- table ketiga hasil emas
tbl_transaction AS (
  SELECT client_id, user_id, mg_date,
  CASE WHEN firm_date IS NULL THEN '2025-03-31' ELSE DATE(firm_date) END date_transaction
  FROM `bigquery-public-data.crypto_transactions`
  LEFT JOIN tbl_firm_mas USING(client_id)
  LEFT JOIN tbl_user_mas USING(client_id)
  WHERE channel_id = 1 AND category = 1 AND status = 3
  AND EXTRACT(YEAR FROM created_at) BETWEEN 2022 AND 2025
  AND DATE(created_at) <= '2025-03-31'
  GROUP BY 1
)
"""

df = client.query(query).to_dataframe()

```

Gambar 3.13 Import Data dari Bigquery
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

```

def get_data(limitation = '2025-03-20', n = 1) label:
    """
    Return first n rows
    ORDER BY sube DESC
    """
    query_sub = client.query(query)
    data = query_sub.to_dataframe()

    # Menampilkan sample data
    print(data.head())

client_id user_id reg_date data_limitation status_type \
0 1815043 921268 2025-03-20 12:27:450000 2025-11-20 1
1 1944201 926053 2025-03-25 18:34:360000 2025-11-23 1
2 2863215 1877025 2025-04-20 16:26:380000 2025-03-21 1
3 2455927 1807598 2025-11-18 18:48:200000 2025-03-21 1
4 1703282 985860 2025-01-20 11:27:480000 2024-07-18 1

last_sub last_reg first_reg first_sub sube \
0 14 60 0 0 976 1.03034e12
1 60 42 118 445 0.180779e11
2 0 26 1 805 4.61094e11
3 20 20 0 766 3.08020e11
4 130 48 889 7.08090e11

rate not_rate rate_avg_rate \
0 7.82259e11 3.33590e11 2136 458 4.58020e03
1 4.28079e11 8.58272e10 540 103 8.30250e03
2 3.82031e11 7.72846e10 4327 1328 1.88542e03
3 7.77324e11 6.96951e10 9640 4326 5.37224e03
4 3.24802e11 2.48320e10 255 460 4.65327e03

avg_rate1_rate1 with_rate1 label \
0 0.0 1000 1
1 0.0 300 1
2 0.0 2070 0
3 1.0 2500 0
4 0.0 824 1

data.info()
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 73225 entries, 0 to 73224

```

Gambar 3.14 Output Import Data dari Bigquery
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Setelah itu praktikan memproses data dengan memisahkan data yang tidak digunakan untuk prediksi (*client_id*) ke dalam *data frame* kosong. *Data frame* ini berfungsi sebagai dasar untuk pemrosesan lebih lanjut, dimana prediksi atau analisis terkait lainnya dapat disimpan untuk setiap klien.

Selanjutnya, mengidentifikasi apakah ada baris yang duplikat dalam dataset. Pemeriksaan ini memastikan bahwa tidak ada data yang terduplikasi yang dimasukkan serta memeriksa apakah ada nilai yang hilang (NaN) dalam dataset. Dengan menghitung jumlah total nilai yang hilang untuk setiap kolom, langkah ini membantu mengidentifikasi masalah potensial dalam data yang perlu ditangani, seperti penghapusan baris yang tidak lengkap, sebelum melanjutkan ke proses analisis atau pembangunan model seperti pada **Gambar 3.15**.

feature engineering

```

[ ] # Create index, label, features dataframe
user_index = data['client_id']
y = data['label']
X = data.drop(columns=['client_id', 'date_last_seen', 'user_id', 'reg_date', 'status_trx',
                        'last_redm', 'subs', 'redm', 'freq_redm', 'with_promo', 'label'])

# scaled_data = StandardScaler().fit_transform(df.copy().values)
scaled_data = RobustScaler().fit_transform(X.values)
scaled_data = pd.DataFrame(data=scaled_data, columns=X.columns.tolist())

scaled_data

```

	last_subs	reg_first_trx	first_last_trx	ret	freq_subs	avg_subs	avg_interval_redm
0	-0.511312	0.076923	2.042289	69079.290101	106.50	236.967245	0.0
1	-0.434389	4.423077	0.736318	18324.121482	28.79	488.458967	0.0
2	-0.523379	-0.115385	1.358209	16000.929997	229.05	48.928533	0.0
3	-0.404721	-0.163848	1.579602	1381.190578	347.70	26.733920	1.0
4	-0.468522	1.364615	1.641791	4987.194097	42.80	198.526334	0.0
...	...	...	...	...	...	...	...
73220	1.227753	0.115385	-0.370647	-0.016958	-0.25	-0.247826	0.0
73221	0.815988	-0.115385	-0.370647	-0.016958	-0.25	-0.247826	0.0
73222	0.702886	10.500000	-0.370647	-0.016958	-0.25	-0.247826	0.0
73223	0.868778	-0.163848	-0.370647	-0.016958	-0.25	-0.247826	0.0
73224	1.110106	0.230769	-0.141791	-0.016047	-0.25	-0.247826	0.0

73225 rows x 7 columns

Gambar 3.16 Feature Engineering Data

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Setelah menyelesaikan *feature engineering* dan normalisasi

- data, perhatian selanjutnya adalah menangani masalah ketidakseimbangan data. Ketidakseimbangan ini akan mempengaruhi kinerja model, terutama ketika satu kelas lebih dominan dibandingkan kelas lainnya. Terdapat tiga metode utama untuk mengatasi masalah ini, yaitu *undersampling*, *oversampling*, dan SMOTE (*Synthetic Minority Over-sampling Technique*). *Undersampling* mengurangi volume data yang tersedia pada kelas mayoritas, sementara *oversampling* meningkatkan volume data pada kelas minoritas. SMOTE meningkatkan keseimbangan data dengan menghasilkan sampel sintesis acak untuk kelas minoritas. Setiap teknik ini diterapkan pada dataset, dan distribusi ketidakseimbangan untuk label setelah penerapan teknik-teknik ini dihitung dan dicetak untuk verifikasi guna memastikan perubahan telah dilakukan seperti pada **Gambar 3.17**.

```
handling imbalance

# Pemisahan features vs target
X = scaled_data

print(X.shape)
print(y.shape)

(73225, 7)
(73225,)

[ ] under = RandomUnderSampler(random_state=42)
over = RandomOverSampler(random_state=42)
smote = SMOTE(random_state=42)

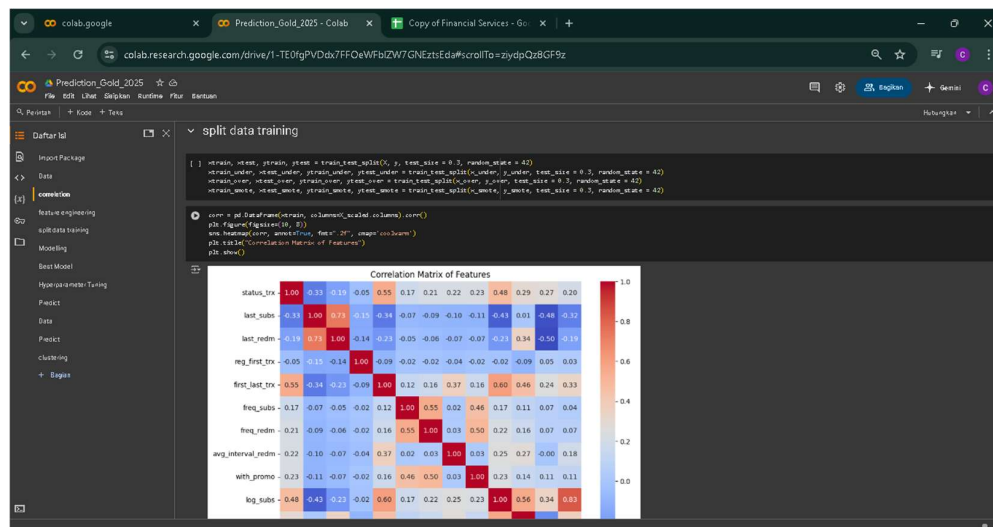
X_under, y_under = under.fit_resample(X, y)
X_over, y_over = over.fit_resample(X, y)
X_smote, y_smote = smote.fit_resample(X, y)

[ ] print('Original')
print(pd.Series(y).value_counts())
print('UNDERSAMPLING')
print(pd.Series(y_under).value_counts())
print('OVERSAMPLING')
print(pd.Series(y_over).value_counts())
print('SMOTE')
print(pd.Series(y_smote).value_counts())

Original
label
0    64971
1     8254
Name: count, dtype: int64
UNDERSAMPLING
label
0     8254
1     8254
Name: count, dtype: int64
OVERSAMPLING
label
0    64971
1     8254
```

Gambar 3.17 Handling Imbalance pada Data
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Selanjutnya, kita perlu membagi data menjadi set pelatihan dan pengujian. Menggunakan fungsi `train_test_split`, data dibagi menjadi dua kelompok dengan alokasi 70% untuk pelatihan dan 30% untuk pengujian. Pembagian ini diterapkan pada empat versi dataset yang telah diproses menggunakan teknik yang berbeda untuk menangani ketidakseimbangan yaitu, dataset asli, dataset yang di-undersample, dataset yang di-oversample, dan dataset yang diproses dengan SMOTE. Hal ini memungkinkan model untuk dilatih dan diuji pada berbagai dataset untuk menilai efektivitas setiap teknik yang digunakan dalam menangani ketidakseimbangan dan meningkatkan hasilnya seperti pada **Gambar 3.18**.



Gambar 3.18 Split Data Menjadi Data Train dan Data Test

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

3. Modelling

Pada bagian ini, Praktikan fokus untuk melatih dan mengevaluasi berbagai model *machine learning* guna menentukan model prediksi terbaik yang dapat memproyeksikan kemungkinan seorang pengguna untuk membeli emas. Langkah pertama adalah menentukan serangkaian *classifier*, yang mencakup *Logistic Regression*, *Random Forest*, *Gradient Boosting*, *CatBoost*, dan *XGBoost* karena mereka banyak digunakan untuk masalah klasifikasi biner. Model-model ini diterapkan pada berbagai versi dataset, yang mencakup dataset asli dan dataset yang telah diproses dengan teknik *undersampling*, *oversampling*, dan *SMOTE* untuk

memperbaiki ketidakseimbangan yang ada dalam data seperti pada **Gambar 3.19**.

```

df_train = pd.DataFrame(columns=['dataset',
                                'precision_train', 'precision',
                                'recall_train', 'recall',
                                'f1_train', 'f1',
                                'accuracy_train', 'accuracy',
                                'auc_train', 'auc'])

classifiers = [
    "Logistic Regression": LogisticRegression(random_state=2),
    "Random Forest": RandomForestClassifier(random_state=2),
    "Gradient Boosting": GradientBoostingClassifier(random_state=2),
    "CatBoost": CatBoostClassifier(random_state=2),
    "XGBoost": XGBClassifier(random_state=2)
]

dataset = [
    'Original': [train_train, test_train],
    'Under Sampling': [train_under, train_under, test_under, test_under],
    'Over Sampling': [train_over, train_over, test_over, test_over],
    'SMOTE': [train_smote, train_smote, test_smote, test_smote]
]

for key in dataset:
    x = dataset[key]
    y = dataset[key]
    x_train = dataset[key][0]
    y_train = dataset[key][1]

    for key in classifiers:
        classifier = classifiers[key]
        model = classifier.fit(x_train, y_train)
        y_pred = model.predict(x_test)
        y_pred_train = model.predict(x_train)

```

Gambar 3.19 Modelling the Data dari Bigquery
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

	precision_train	precision	recall_train	recall	f1_train	f1	auc_train	auc	accuracy_train
0	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
1	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
2	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
3	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
4	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
5	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
6	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
7	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
8	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
9	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
10	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
11	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
12	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
13	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
14	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
15	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
16	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
17	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
18	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
19	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082
20	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082	0.963082

Gambar 3.20 Output Modelling Data dari Bigquery
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

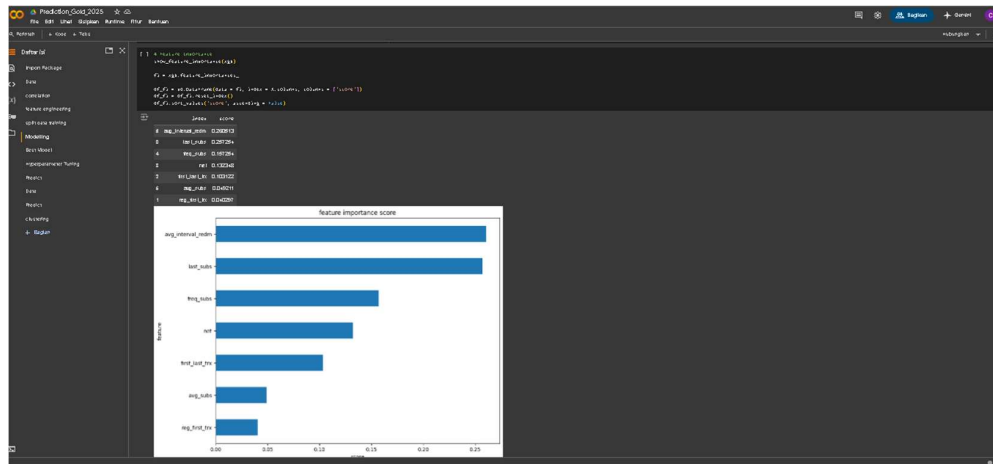
Setiap *classifier* dilatih pada data latih dan dievaluasi pada data uji yang sesuai. Berbagai metrik evaluasi, termasuk *precision*, *recall*, *F1-score*, akurasi, AUC, digunakan untuk menilai performa

model, baik pada fase pelatihan maupun pengujian. Metrik hasil ini mendefinisikan evaluasi terhadap masing-masing model dan disimpan dalam *data frame* (df_eval) untuk analisis lebih lanjut. Prediksi model dibandingkan dengan hasil aktual, sehingga dapat menentukan kredibilitas dan validitas model dalam meramalkan hasil yang diinginkan seperti pada **Gambar 3.21**.

	model	dataset	precision_train	precision	recall_train	recall	f1_train	f1	accuracy_train	accuracy	auc_train	auc
11	RandomForest	Over Sampling	1.00	0.90	1.00	1.00	1.00	0.98	1.00	0.98	1.00	1.00
16	RandomForest	SMOTE	1.00	0.93	1.00	0.95	1.00	0.94	1.00	0.94	1.00	0.98
10	CarBoost	SMOTE	0.98	0.95	0.95	0.93	0.96	0.94	0.96	0.94	0.99	0.99
13	CarBoost	Over Sampling	0.91	0.89	0.94	0.92	0.92	0.90	0.92	0.90	0.97	0.97
14	XGBoost	Over Sampling	0.91	0.89	0.94	0.92	0.92	0.91	0.92	0.90	0.98	0.97
19	XGBoost	SMOTE	0.94	0.93	0.94	0.92	0.94	0.92	0.94	0.93	0.99	0.98
17	GradientBoosting	SMOTE	0.87	0.86	0.90	0.89	0.88	0.88	0.88	0.88	0.95	0.95
12	GradientBoosting	Over Sampling	0.84	0.84	0.86	0.86	0.85	0.85	0.85	0.85	0.93	0.93
9	XGBoost	Under Sampling	0.93	0.84	0.94	0.85	0.94	0.84	0.94	0.85	0.99	0.93
8	CarBoost	Under Sampling	0.90	0.85	0.90	0.85	0.90	0.85	0.90	0.85	0.97	0.93
7	GradientBoosting	Under Sampling	0.84	0.82	0.86	0.85	0.85	0.83	0.85	0.83	0.93	0.92
15	LogisticRegression	SMOTE	0.78	0.75	0.85	0.85	0.80	0.80	0.79	0.79	0.87	0.87
10	LogisticRegression	Over Sampling	0.75	0.75	0.84	0.84	0.80	0.79	0.79	0.78	0.86	0.87
6	RandomForest	Under Sampling	1.00	0.85	1.00	0.84	1.00	0.84	1.00	0.84	1.00	0.92
5	LogisticRegression	Under Sampling	0.76	0.75	0.84	0.83	0.80	0.79	0.79	0.78	0.87	0.86
4	XGBoost	Original	0.88	0.77	0.95	0.54	0.75	0.64	0.95	0.92	0.97	0.93
3	CarBoost	Original	0.88	0.79	0.93	0.53	0.73	0.64	0.95	0.92	0.96	0.94
1	RandomForest	Original	1.00	0.79	1.00	0.51	1.00	0.62	1.00	0.92	1.00	0.92
2	GradientBoosting	Original	0.82	0.82	0.46	0.47	0.60	0.60	0.93	0.93	0.93	0.92
0	LogisticRegression	Original	0.81	0.83	0.69	0.69	0.76	0.76	0.89	0.90	0.88	0.86

Gambar 3.21 Classification Report Modelling Berdasarkan Model Terbaik

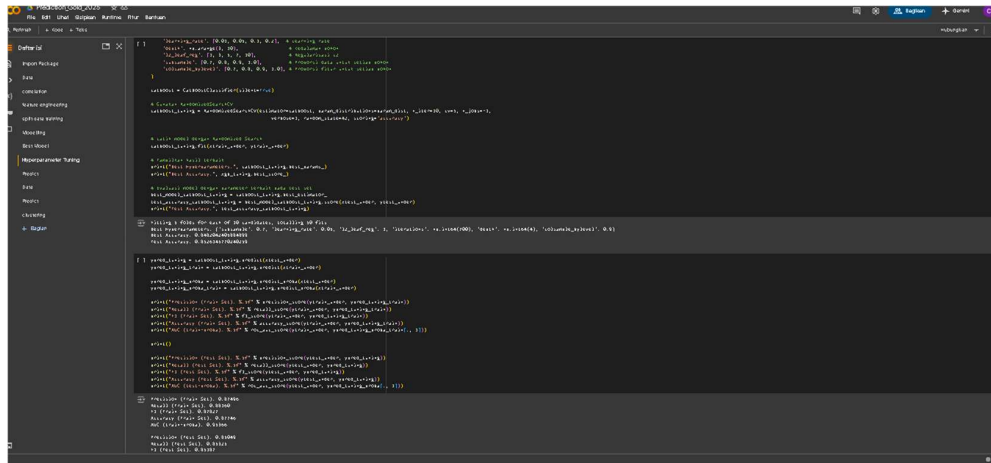
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan



Gambar 3.22 Mencari *Feature Importance* pada Model Terbaik

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Berdasarkan *classification reports*, model yang terpilih adalah *XGBoost* dan *Catboost* dengan teknik *undersampling* karena dinilai lebih stabil dan tidak *overfitting*. Kemudian, kedua model tersebut dicari *feature importancenya* dan dilakukan tuning dengan optimasi *hyperparameter* menggunakan *RandomizedSearchCV*. Proses ini mengoptimalkan model dengan parameter kunci seperti learning rate, jumlah estimator, kedalaman pohon, dan rasio subsample pada model *XGBoost* dan *CatBoost*. Model yang menghasilkan akurasi tertinggi selama pencarian dipilih untuk optimasi *hyperparameter* lebih lanjut seperti pada **Gambar 3.22** dan **Gambar 3.33**



Gambar 3.23 Melakukan *Hyperparameter Tuning* pada Model Terbaik
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Model yang telah di-*tuning hyperparameternya* kemudian dievaluasi kembali menggunakan data uji dan menghasilkan laporan klasifikasi yang komprehensif dengan *precision*, *recall*, *F1-score*, akurasi, dan AUC. Selain itu, sebuah kurva *precision-recall* juga dipetakan untuk membandingkan *trade-off* antara *recall* dan *precision* pada berbagai level *threshold*. Grafik ini juga menampilkan *feature importance* dari model, yang menunjukkan fitur mana yang paling berpengaruh pada hasil prediksi. Prosedur ini memastikan bahwa model disempurnakan untuk memberikan prediksi yang paling akurat mengenai pembelian emas di masa depan. Hasil model yang telah di-*tuning* kemudian disimpan agar dapat digunakan lagi pada data prediksi seperti pada **Gambar 3.24**.

```

~ Predict

[ ] import pickle
    from google.colab import drive

    # Menghubungkan Google Drive ke Google Colab
    drive.mount('/content/drive')

    # Tentukan path file model yang sudah disimpan
    xgb_model_path = '/content/drive/MyDrive/xgb_tuning.pkl' # Path ke file model XGBoost
    catboost_model_path = '/content/drive/MyDrive/catboost_tuning.pkl' # Path ke file model CatBoost

    # Memuat model XGBoost
    with open(xgb_model_path, 'rb') as file:
        xgb_tuning = pickle.load(file)

    # Memuat model CatBoost
    with open(catboost_model_path, 'rb') as file:
        catboost_tuning = pickle.load(file)

    # Menampilkan pesan bahwa model berhasil dimuat
    print("Model XGBoost berhasil dimuat.")
    print("Model CatBoost berhasil dimuat.")

Drive already mounted at /content/drive; to attempt to forcibly remount, call drive.mount("/content/drive", force_remount=True).
Model XGBoost berhasil dimuat.
Model CatBoost berhasil dimuat.

```

Gambar 3.24 Menyimpan Model
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

4. Predict Data

Ketika model dianggap siap, tugas selanjutnya adalah memprediksi data baru menggunakan model ini. Langkah pertama melibatkan persiapan data yang akan diprediksi, termasuk menerapkan *feature importance* yang sama dengan yang diterapkan pada data pelatihan. Langkah ini diperlukan agar fitur yang ada pada data baru dan yang digunakan pada data pelatihan relevan dan bermakna seperti pada **Gambar 3.25** dan **Gambar 3.26**.

```

Predict

~ Data

Query #1

-- Select output table name
FROM `bigquery-public-data`.`ml`.`xgb_tuning` AS xgb_tuning
SELECT CASE WHEN (xgb_tuning['xgb_tuning'] IS NOT NULL) THEN 'xgb_tuning' ELSE 'catboost_tuning' END AS output_table_name
FROM xgb_tuning
WHERE (xgb_tuning['xgb_tuning'] IS NOT NULL)
ORDER BY 1

Query #2

-- Select output table name
FROM `bigquery-public-data`.`ml`.`catboost_tuning` AS catboost_tuning
SELECT CASE WHEN (catboost_tuning['catboost_tuning'] IS NOT NULL) THEN 'catboost_tuning' ELSE 'xgb_tuning' END AS output_table_name
FROM catboost_tuning
WHERE (catboost_tuning['catboost_tuning'] IS NOT NULL)
ORDER BY 1

Query #3

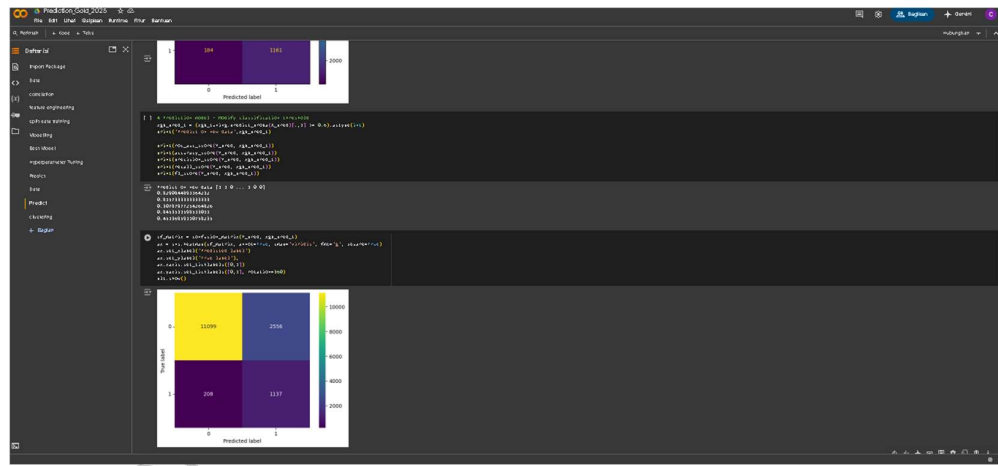
-- Select output table name
FROM `bigquery-public-data`.`ml`.`xgb_tuning` AS xgb_tuning
SELECT CASE WHEN (xgb_tuning['xgb_tuning'] IS NOT NULL) THEN 'xgb_tuning' ELSE 'catboost_tuning' END AS output_table_name
FROM xgb_tuning
WHERE (xgb_tuning['xgb_tuning'] IS NOT NULL)
ORDER BY 1

Query #4

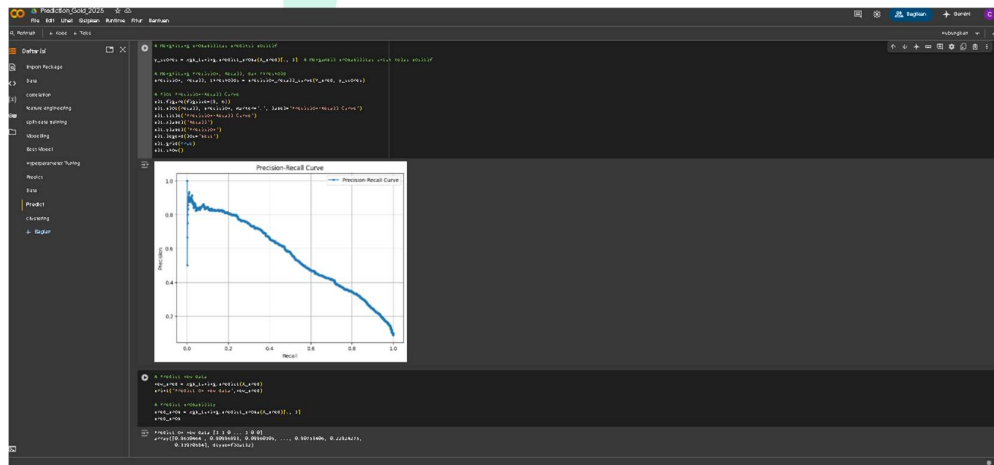
-- Select output table name
FROM `bigquery-public-data`.`ml`.`catboost_tuning` AS catboost_tuning
SELECT CASE WHEN (catboost_tuning['catboost_tuning'] IS NOT NULL) THEN 'catboost_tuning' ELSE 'xgb_tuning' END AS output_table_name
FROM catboost_tuning
WHERE (catboost_tuning['catboost_tuning'] IS NOT NULL)
ORDER BY 1

```

Gambar 3.25 Import Data Baru dari BigQuery
 Sumber : Dokumentasi yang Dihasilkan oleh Praktikan



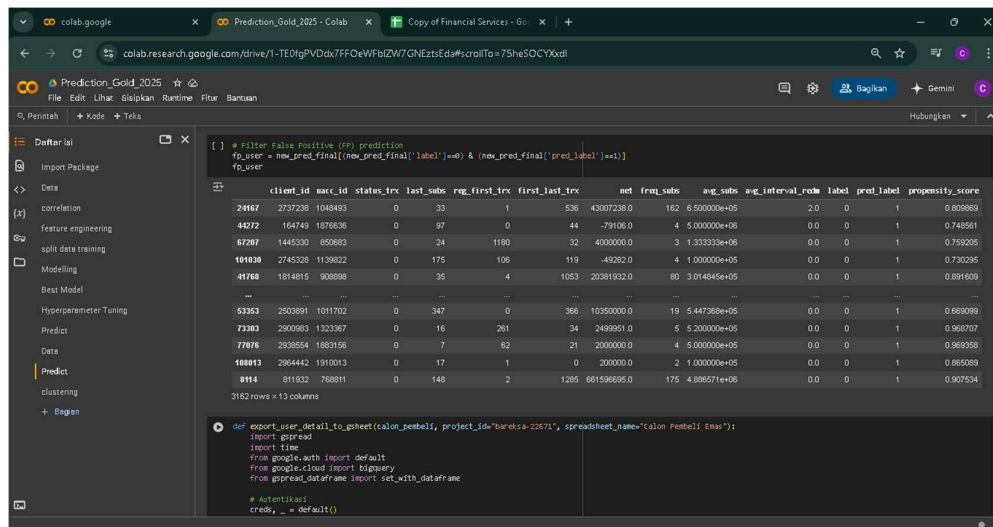
Gambar 3.28 Training Model pada Data Baru (a)
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan



Gambar 3.29 Training Model pada Data Baru (b)
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Hasil dari prediksi dievaluasi (dinilai), salah satu langkah utama adalah deteksi *false positives* (FP) yang menunjukkan pengguna yang diprediksi tidak akan membeli emas, tetapi sebenarnya mereka diprediksi akan membeli produk tersebut. Mendeteksi *false positives* berguna karena membantu memberikan panduan untuk rencana pemasaran yang lebih baik di masa depan.

Dengan kata lain, *false positives* disini adalah pengguna yang salah diasumsikan oleh model tidak menunjukkan minat untuk membeli emas tapi ternyata tertarik pada produk emas, sehingga tindakan bisa diambil untuk menghindari pemborosan sumber daya dan upaya seperti pada **Gambar 3.30**.



```

[ ] # Filter False Positive (FP) prediction
fp_user = new_pred_final[(new_pred_final['label']==0) & (new_pred_final['propensity_score']>0)]
fp_user

```

client_id	name	email	phone	propensity_score
24167	2737238	1048493	0	33
44272	164749	1876636	0	97
67207	1445330	850683	0	24
101030	2745328	1138222	0	175
41768	1814815	908898	0	35
53353	2503891	1011702	0	347
73303	2900983	1323367	0	16
77076	2938554	1883156	0	7
108013	2964442	1910013	0	17
8114	811832	760811	0	148

```

def export_user_detail_to_gsheet(calon_pembeli, project_id="bareksa-22671", spreadsheet_name="Calon Pembeli Emas"):
    import gspread
    import time
    from google.auth import default
    from google.cloud import bigquery
    from gspread_dataframe import set_with_dataframe

    # Autentikasi
    creds, _ = default()

```

Gambar 3.30 Filter False Positive User di Bigquery
Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Setelah melakukan prediksi dan mengevaluasi dimana *false positives* terjadi, informasi yang relevan diekspor ke *Google Sheets* untuk analisis lebih dalam dan upaya kolaboratif. Informasi yang diekspor mencakup data pengguna seperti *client_id*, nama, email, dan nomor telfon. Data ini difilter sesuai dengan *propensity score* yang mengukur kemungkinan seorang pengguna membeli emas berdasarkan perhitungan model. Mekanisme ekspor ke *Google Sheets* memungkinkan tim lain seperti pemasaran untuk mengakses dan mengambil tindakan terhadap pengguna yang diprediksi sebagai calon pembeli emas, memfokuskan upaya pemasaran mereka pada pengguna dengan skor *propensity* tertinggi seperti pada **Gambar 3.31**.

id	nama	alamat	telepon	email	password
1	John Doe	123 Main St	555-123-4567	john.doe@gmail.com	123456
2	Jane Smith	456 Oak St	555-987-6543	jane.smith@gmail.com	654321
3	Bob Johnson	789 Pine St	555-234-5678	bob.johnson@gmail.com	876543
4	Alice Brown	101 Elm St	555-345-6789	alice.brown@gmail.com	987654
5	Charlie Davis	202 Maple St	555-456-7890	charlie.davis@gmail.com	098765
6	Diana Prince	303 Cedar St	555-567-8901	diana.prince@gmail.com	109876
7	Frank Miller	404 Birch St	555-678-9012	frank.miller@gmail.com	210987
8	Grace Wilson	505 Spruce St	555-789-0123	grace.wilson@gmail.com	321098
9	Henry Taylor	606 Willow St	555-890-1234	henry.taylor@gmail.com	432109
10	Ivy Clark	707 Ash St	555-901-2345	ivy.clark@gmail.com	543210

Gambar 3.31 Output di Google Sheet

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

3.2.4 Visualisasi Data

1. Gold Weekly New Registered

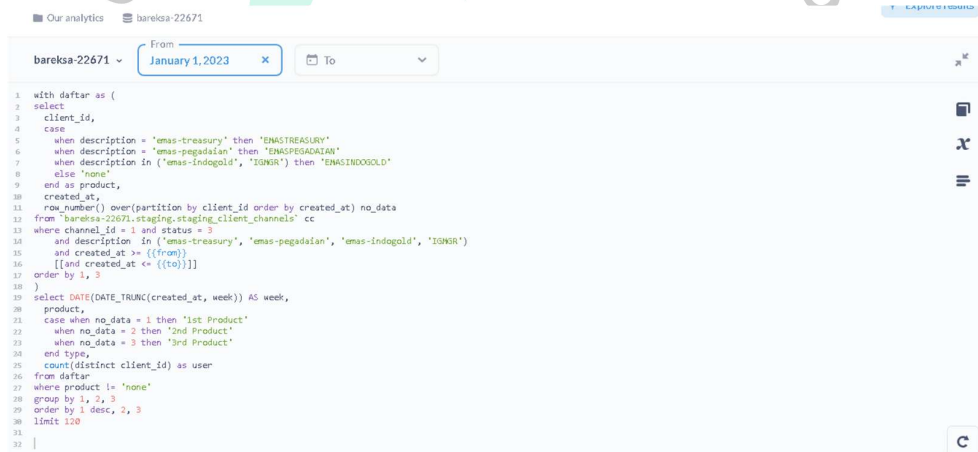
Visualisasi data ini bertujuan untuk melacak pengguna yang baru registrasi pada setiap produk emas secara *real-time*. Dalam kueri ini, deskripsi produk seperti emas-*treasury*, emas-pedagaaian, dan emas-indogold dipetakan ke pengenalan kategori tertentu yaitu EMASTREASURY, EMASPEDAGAIAN, dan EMASINDOGOLD. Jika ada deskripsi produk yang tidak cocok dengan ini, maka akan diklasifikasikan sebagai 'none.' Klasifikasi ini membantu untuk mengetahui produk mana yang menarik perhatian pengguna untuk keterlibatan lebih lanjut.

Selain itu, dalam *query* ini, kategori pengguna diterapkan dengan menggunakan fungsi `row_number()`, yang dipartisi berdasarkan `client_id` dan diurutkan berdasarkan tanggal `created_at`, untuk melacak transaksi pengguna secara kronologis. Selanjutnya, *query* ini membatasi data hanya untuk mencakup catatan yang memenuhi kriteria `channel_id = 1` yang mendeklarasikan produk emas dan `status = 3` yang mendeklarasikan pengguna berhasil terdaftar. Selain itu, tanggal transaksi *difilter* untuk memastikan bahwa hanya transaksi yang terjadi setelah tanggal tertentu `{{from}}`

yang diambil, yang dimaksudkan agar visualisasi data ini fokus pada aktivitas terbaru.

Pada bagian kedua *query*, jumlah pengguna yang terdaftar untuk setiap jenis produk dihitung dan dikelompokkan berdasarkan jenis produk (product). Pengguna diklasifikasikan sebagai terdaftar untuk *1st Product*, *2nd Product*, atau *3rd Product*, tergantung pada minat mereka terhadap produk tersebut, yang membantu untuk melihat produk mana yang lebih banyak diminati oleh pengguna baru. Kueri ini diurutkan berdasarkan tanggal pendaftaran produk untuk memprioritaskan pendaftaran terbaru, dan hasilnya dibatasi hingga 120 catatan untuk memfokuskan hasil dataset seperti pada

Gambar 3.32.



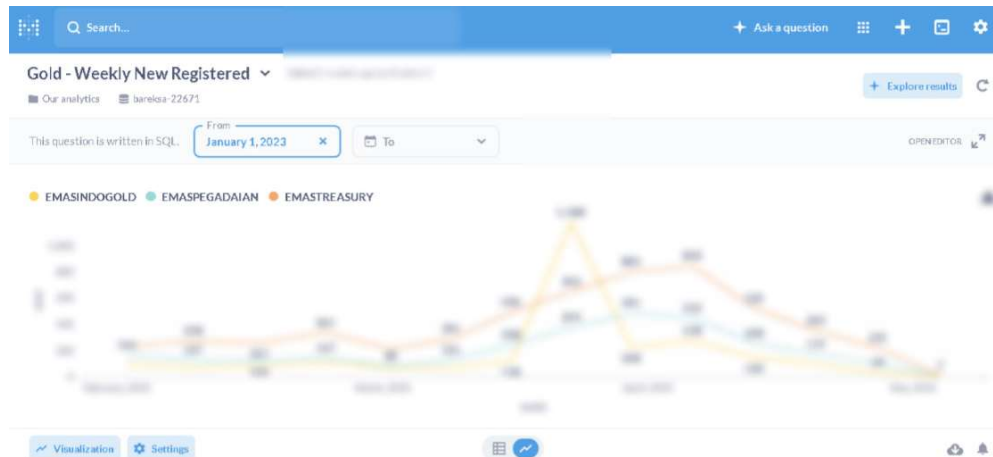
```
1 with daftar as (  
2   select  
3     client_id,  
4     case  
5       when description = 'emas-treasury' then 'ENASTREASURY'  
6       when description = 'emas-pegadaian' then 'ENASPEGADAIAN'  
7       when description in ('emas-indogold', 'IGNWR') then 'ENASINDOGOLD'  
8       else 'none'  
9     end as product,  
10    created_at,  
11    row_number() over(partition by client_id order by created_at) no_data  
12  from "bareksa-22671.staging.staging_client_channels" cc  
13  where channel_id = 1 and status = 3  
14    and description in ('emas-treasury', 'emas-pegadaian', 'emas-indogold', 'IGNWR')  
15    and created_at >= ((from))  
16    [[and created_at <= ((to))]]  
17  order by 1, 3  
18 )  
19 select DATE_TRUNC(created_at, week) AS week,  
20 product,  
21 case when no_data = 1 then '1st Product'  
22    when no_data = 2 then '2nd Product'  
23    when no_data = 3 then '3rd Product'  
24 end type,  
25 count(distinct client_id) as user  
26 from daftar  
27 where product != 'none'  
28 group by 1, 2, 3  
29 order by 1 desc, 2, 3  
30 limit 120  
31  
32 |
```

Gambar 3.32 Proses Query di Metabase

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Analisis ini memberikan gambaran mengenai tren pendaftaran pengguna, distribusi minat terhadap produk yang berbeda, dan tingkat aktivitas pengguna baru dalam periode waktu tertentu. Bisnis dapat melacak pengguna berdasarkan aktivitas pendaftaran dan minat produk, yang memberikan pemahaman yang lebih baik tentang interaksi pelanggan, mendeteksi peluang baru, dan memungkinkan perusahaan untuk mengembangkan rencana

yang lebih terarah untuk menarik perhatian pelanggan potensial seperti pada **Gambar 3.33**



Gambar 3.33 Output Weekly Registered di Metabase

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

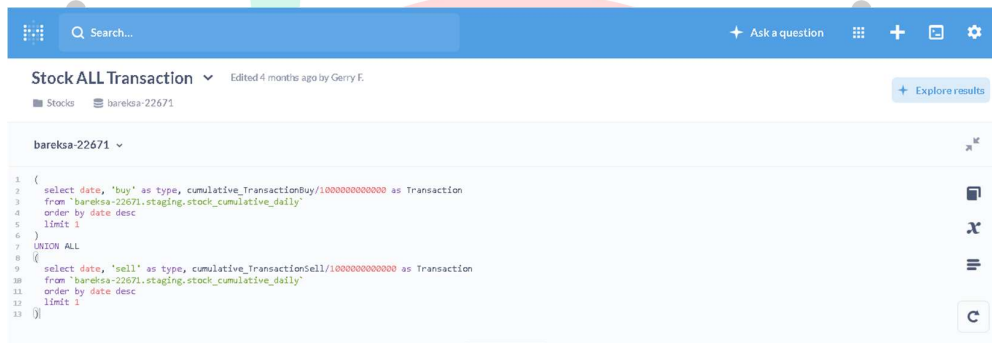
2. Transaksi Stock

Visualisasi data transaksi *stock* bertujuan untuk melacak transaksi penjualan dan pembelian dalam saham secara *real-time*.

Query mengambil data transaksi beli dan jual saham terbaru dari tabel *stock_cumulative_daily*, yang berisi transaksi saham harian secara kumulatif. *Query* ini menggunakan operator *UNION ALL* untuk menggabungkan kedua set transaksi beli dan jual. Pada bagian pertama, *query* mengambil data yang berkaitan dengan transaksi pembelian saham. Kolom *cumulative_TransactionBuy* diambil dari tabel *stock_cumulative_daily* dan untuk memudahkan pemahaman, nilai transaksi pembelian kumulatif yang diambil dibagi dengan 1 triliun (1000000000000). Kolom *date* memberikan informasi terkait tanggal transaksi, sementara kolom *type* diisi dengan *buy*, yang berarti transaksi ini adalah pembelian. Hasilnya kemudian diurutkan berdasarkan kolom *date* dalam urutan menurun, sehingga transaksi terbaru akan muncul pertama kali. *Query* ini

menggunakan klausa *LIMIT* 1 untuk memastikan hanya transaksi pembelian terbaru yang diambil.

Bagian kedua dari *query* menjalankan operasi yang serupa, namun kali ini untuk transaksi penjualan saham. Data yang diambil berasal dari kolom *cumulative_TransactionSell*, yang berisi jumlah kumulatif nilai transaksi penjualan saham. Nilai di kolom ini juga dibagi dengan satu triliun untuk memudahkan pemahaman. Kolom *date* menunjukkan tanggal transaksi, sementara kolom *type* diisi dengan nilai *sell*, yang berarti transaksi ini mencakup penjualan. Seperti pada bagian pertama, data diurutkan berdasarkan kolom tanggal dalam urutan menurun, dan hanya satu baris terbaru yang diambil menggunakan perintah *LIMIT* 1 seperti pada **Gambar 3.34**.

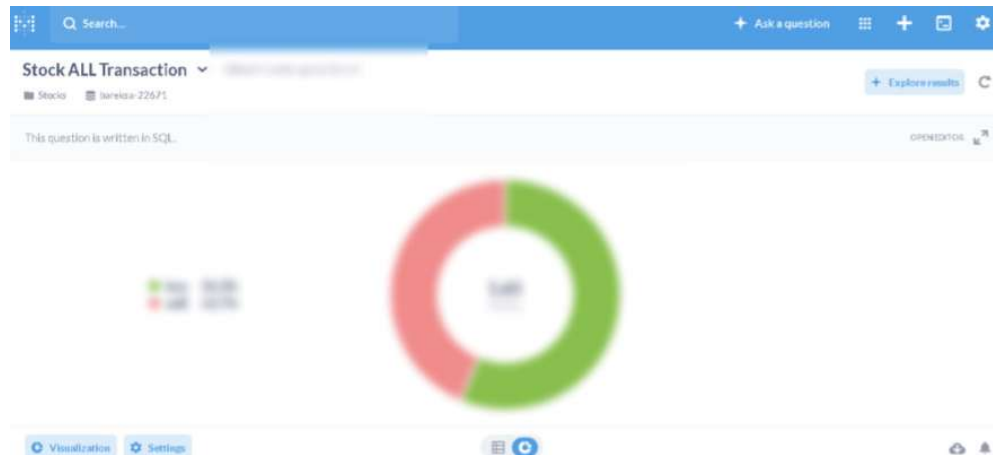
The image shows a screenshot of the Metabase query editor interface. At the top, there is a search bar and a navigation menu. Below the search bar, the query is titled "Stock ALL Transaction" and is attributed to "Gerry F.". The query is written in SQL and uses the *UNION ALL* operator to combine two queries. The first query selects the date, type (buy), and cumulative transaction buy amount (divided by 1,000,000,000,000) from the *bareksa-22671.staging.stock_cumulative_daily* table, ordered by date descending and limited to 1 row. The second query selects the date, type (sell), and cumulative transaction sell amount (divided by 1,000,000,000,000) from the same table, also ordered by date descending and limited to 1 row. The interface includes a "Search..." bar, a "Ask a question" button, and a "Explore results" button. The query is displayed in a code editor with line numbers on the left and a "Copy" button on the right.

Gambar 3.34 Proses Query Stock di Metabase

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

Kedua set data ini digabungkan menjadi satu hasil menggunakan klausa *UNION ALL*. Dengan menggunakan *UNION ALL*, hasil transaksi beli dan transaksi jual akan disertakan tanpa menghapus baris yang duplikat. Hasil akhirnya adalah dua baris; satu berisi transaksi pembelian terbaru dengan tipe *buy*, dan satu lagi berisi transaksi penjualan terbaru dengan tipe *sell*. Dengan cara ini, jumlah transaksi pembelian dan penjualan saham yang tersedia dapat ditangkap secara *real-time* dan dapat digunakan untuk analisis

lebih lanjut atau pelaporan terkait dengan aktivitas pasar saham seperti terlihat pada **Gambar 3.35**.



Gambar 3.35 Output Stock Transaction pada Metabase

Sumber : Dokumentasi yang Dihasilkan oleh Praktikan

3.3 Kendala Yang Dihadapi

Selama menjalankan tugas kerja praktik, praktikan menghadapi beberapa kendala yang cukup menghambat kelancaran pekerjaan diantaranya adalah

1. Ketidapahaman atas *database* yang dimiliki perusahaan sehingga pada saat proses pengumpulan data mentah yang dilakukan secara mandiri tanpa adanya penjelasan yang jelas mengenai makna dan konteks data tersebut, membuat praktikan kesulitan dalam menginterpretasikan data secara tepat dan mengarahkannya pada analisis yang relevan.
2. Miskomunikasi antar divisi, yang mengakibatkan duplikasi dalam pekerjaan dan memperlambat progres proyek. Misalnya, beberapa tim yang bekerja pada tugas yang sama tidak saling berkomunikasi

dengan baik, sehingga satu pekerjaan bisa dikerjakan oleh dua divisi tanpa sepengetahuan satu sama lain. Hal ini menyebabkan pemborosan waktu dan sumber daya yang seharusnya dapat dialokasikan untuk tugas lain.

3. Data mengalami perbaikan atau perubahan yang dilakukan oleh tim *backend engineer*. Ketidakpastian data mengganggu kestabilan hasil analisis yang telah dilakukan sebelumnya, sehingga praktikan harus melakukan penyesuaian dan perbaikan yang memakan waktu ekstra. Kesulitan-kesulitan ini memperlambat proses dokumentasi dan pemrosesan data yang sudah terkumpul.

3.4 Cara Mengatasi Kendala

Untuk mengatasi berbagai kendala yang dihadapi selama masa kerja praktik, praktikan mengambil beberapa langkah yang dinilai efektif untuk mengurangi dampak dari masalah-masalah tersebut yaitu

1. Praktikan melakukan diskusi lebih lanjut dengan mentor untuk membahas klasifikasi dan pemahaman tentang data yang digunakan. Ini membantu praktikan dalam mengklarifikasi makna dan tujuan dari data yang dikumpulkan, serta memastikan bahwa pengolahan data dilakukan dengan cara yang benar dan efektif.
2. Mengenai masalah miskomunikasi antar divisi, praktikan secara proaktif melakukan konfirmasi dengan mentor terkait *request* atau permintaan dari divisi lain yang langsung diteruskan kepada praktikan. Ini membantu mengurangi risiko kesalahpahaman

dan memastikan bahwa setiap permintaan yang diterima sudah jelas dan terarah.

3. Pada masalah perbaikan data dari tim *backend engineer*, praktikan mengkonfirmasi dengan mentor untuk mendapatkan kejelasan data yang digunakan, serta menunggu validasi data yang diperlukan untuk analisis lebih lanjut serta menetapkan batas waktu pengumpulan data agar analisis dilakukan pada data yang sudah terkumpul secara final.

3.5 Pembelajaran Yang Diperoleh dari Kerja Profesi

Selama menjalankan kerja praktik di PT Bareksa, praktikan memperoleh berbagai pembelajaran yang sangat berharga, baik dalam hal keterampilan teknis maupun komunikasi. Salah satu pembelajaran utama yang diperoleh adalah peningkatan keterampilan dalam analisis data. Praktikan belajar untuk mengelola data mentah yang besar dan kompleks menggunakan berbagai *tools* dan bahasa pemrograman seperti *Python*, *BigQuery*, dan *Metabase*. Praktikan juga menjadi lebih terampil dalam memvisualisasikan data yang kompleks sehingga dapat disajikan dengan cara yang mudah dipahami oleh berbagai pihak, termasuk manajer produk dan tim pemasaran. Selain itu, praktikan memperoleh pemahaman yang lebih dalam mengenai produk investasi dan bagaimana peran data sangat krusial dalam pengembangan produk dan strategi bisnis perusahaan. Praktikan juga memahami bagaimana teknik *machine learning* dapat digunakan untuk memprediksi perilaku pengguna dan meningkatkan efektivitas pemasaran produk. Serta dapat menerapkan bidang ilmu yang telah dipelajari pada Mata Kuliah *Business Intelligence* dan *Data Warehouse*.

Selain keterampilan teknis, praktikan juga mengembangkan keterampilan komunikasi yang sangat penting dalam dunia kerja. Praktikan belajar untuk berkomunikasi dengan jelas dan tepat

waktu dengan berbagai divisi dalam perusahaan untuk memastikan bahwa setiap pekerjaan berjalan dengan baik dan sesuai dengan tujuan yang telah ditetapkan. Praktikan juga melakukan eksplorasi terhadap tools yang sebelumnya belum digunakan, seperti *BigQuery* dan *Metabase*, yang meningkatkan kemampuan praktikan dalam menangani data skala besar dan memberikan *insight* yang lebih mendalam. Secara keseluruhan, pengalaman kerja praktik ini memberikan pembelajaran yang sangat berharga yang dapat diterapkan dalam dunia kerja yang lebih profesional dan teknis.

