

SAP Data Audit Masterclass

Python cheat sheet

All the Python you need for all the 300 must have data audit tests

80% of what you need to know for Python for the 300 must-have data analytics tests in the 300Framework Masterclass

Legend

Tip

Note: all of the - that you see in the code should be the short - (minus) and not the long – used for writing.

Links will be shown in blue font

Code in the CMD box

Code in a PowerShell

Code in Visual Studio Code PowerShell terminal

Code in a Python script

Code in other file types: .toml, .txt, .env, .json

Set-up

in these pages we go through the pre-requisites that you need, that will open up the Python world and all its possibilities to you.

After completing set-up, you will be able to use python to do pretty much anything (*within the same access limits that you have*), from reading images, finding information on the internet, sending emails, analyzing data, using artificial intelligence, gathering data from databases, files and systems, sending results to PBI or other systems, creating an interactive AI intranet site and even making your PC talk to you....

Set-up

To start using Python, you need to install Python on your machine, as well as an IDE (Integrated Development Environment), which is basically like a glorified text editor that helps you to edit Python code and run programs.

1/ Download the version of Python that you require

To get set-up with Python you need to first download Python and install it on your computer:

<https://www.python.org/downloads>



In this document we will assume that we are using Python version 3.14. However, if you are reading this document at a later date in the future, you will probably see a later version, than that indicated in the screen/shot above.

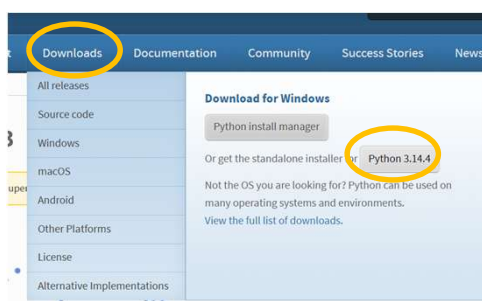
2/ (Alternative) Download an earlier version of Python

You can choose to install a version of Python that is earlier than the latest one, by scrolling down the screen of the above site and choosing the version that you need:

Looking for a specific release?			
Python releases by version number:			
Release version	Release date		Click for more
Python 3.14.4	April 7, 2026	Download	Release notes
Python 3.14.3	Feb. 3, 2026	Download	Release notes

For example, if you want to run a project that you created in the past, with Python library versions that are from the past, then you would want to use the Python version that you used in the past. Otherwise, you may find that you have conflicts between the versions of Python libraries that you will add to your project and your newer Python version.

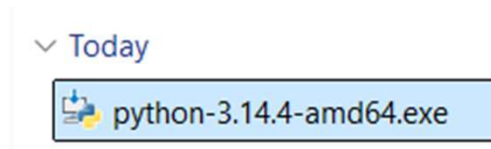
After you press one of the download buttons from these earlier Python releases, you will be redirected to a page about that version. Click on the button Downloads at the top of this page and then click on the button with the version number to get the standalone executable file. You should see that a .exe file is downloaded to your computer.



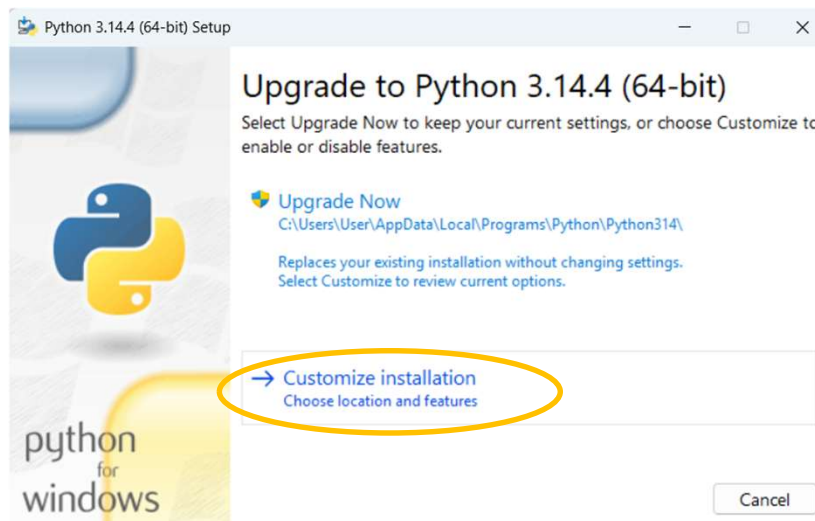
Set-up

3/ Install the downloaded Python program

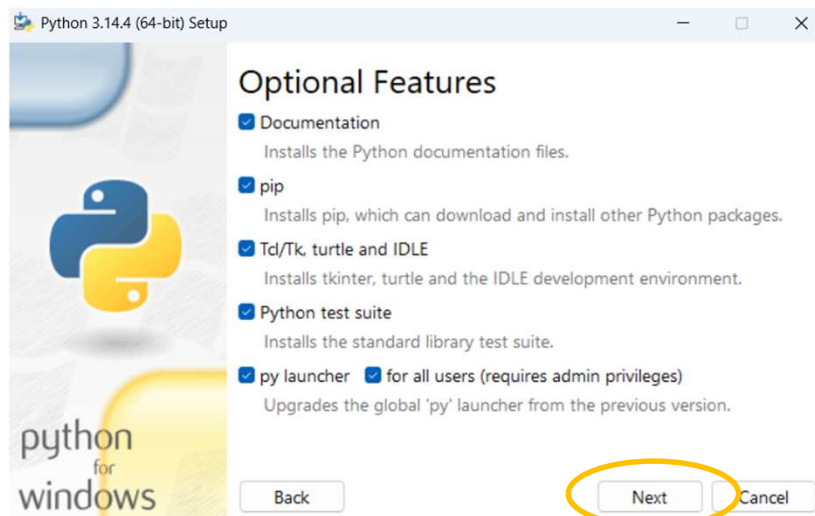
Double-click on the .exe that you downloaded to launch the installer.



Choose Customize installation:

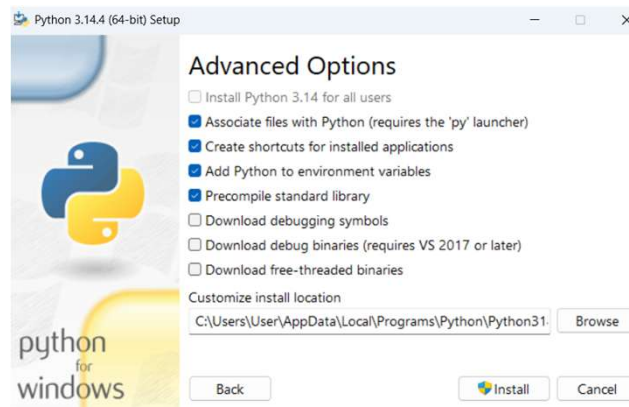


Keep all the optional features checked, then click next:

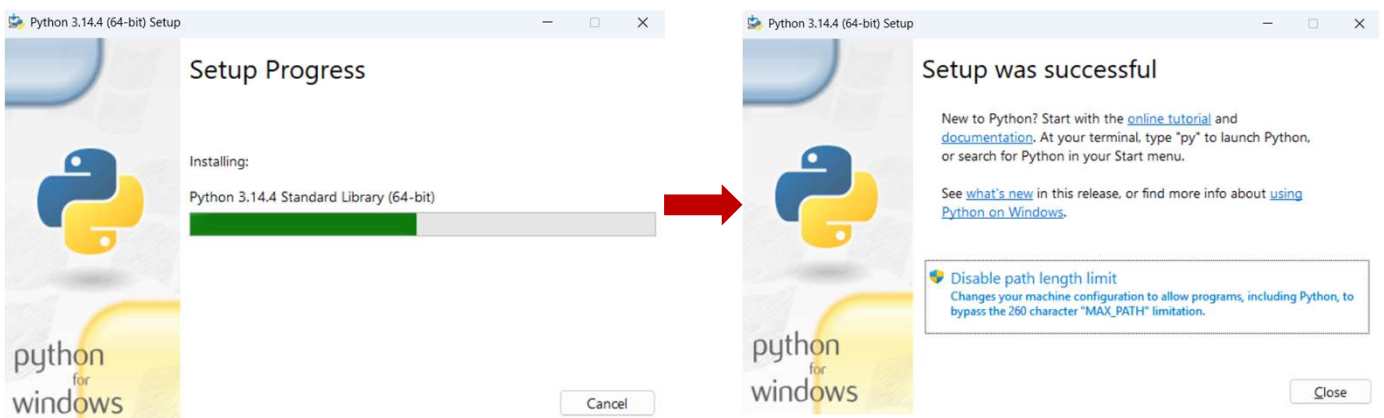


Set-up

in advanced options, ensure that **Add Python to environment variables** (this ensures that Python can be found by your computer) and **Precompile standard library** are checked:



You can then click on Install and then close.



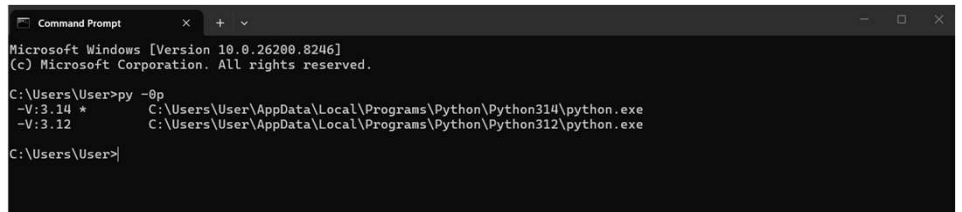
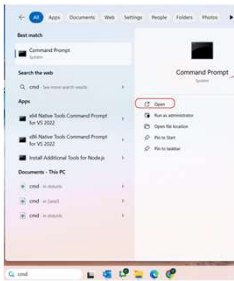
Set-up

4/ Check that the Python version installed correctly

Following installation of Python, type CMD in the Windows Search bar to open a command Window, at the default location of your PC. Then type the following:

```
py -0p
```

You will then see all of the Python versions that exist on your PC. You can have more than one Python version on your PC. When we create a project, we will create a virtual environment for our project that will use one of these specific Python versions. You will only be able to create a virtual environment that has a Python version that is found on your machine.



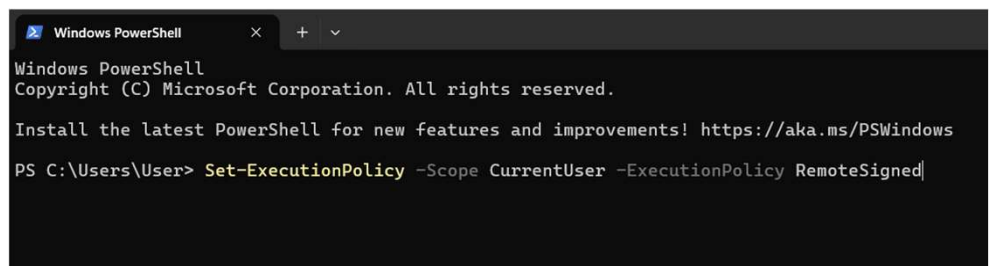
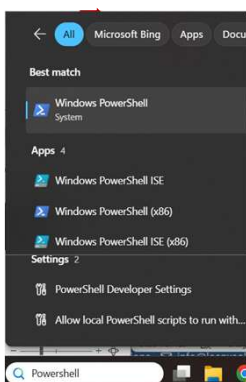
Note: It is important to note the Python version used in any project. This way, if you want to re-use the same project in the future, you can download and install the correct Python version for that project. This way, you can avoid any incompatibility bugs between Python version and library versions. In the section on Project Organization, we provide best practices to ensure that we always document correctly the Python version and library versions used in our projects.

5/ (If necessary for your PC) Allow yourself to run programs

When you get to the Python steps, in the following pages, you may find that your PC does not allow you to run programs. If this is the case, or to avoid any potential future issues, you can set your PC to allow you to run programs.

We will set our PC to be able to run programs, by using a PowerShell command. To open a PowerShell on Windows, type PowerShell in the Search bar. Then type the following command in the Window that opens and click enter:

```
Set-ExecutionPolicy -Scope CurrentUser -ExecutionPolicy RemoteSigned
```



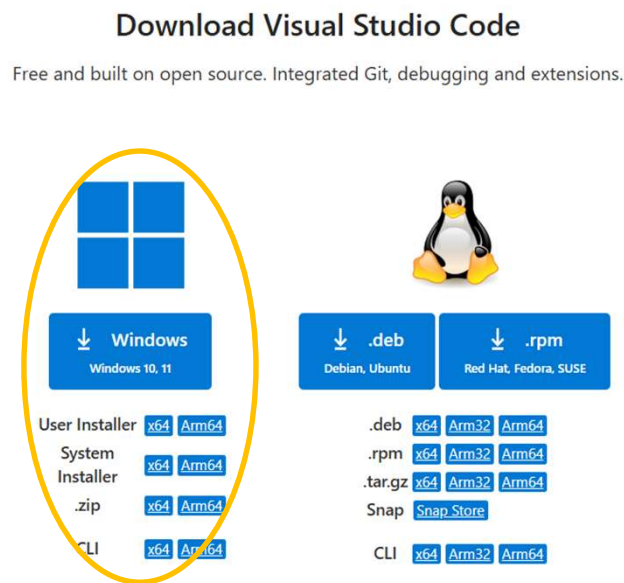
Set-up

6/ Download the Integrated Development Environment (IDE), Visual Studio Code (VSC)

We will use the free IDE, Visual Studio Code (VSC). We use VSC because it is the simplest IDE to use. VSC colors the code that we write (so that we can read it quicker). VSC also has a terminal window, which enables us to run CMD or PowerShell commands, as we have been doing in the previous steps. From VSC, we can therefore also run the Python programs that we will create.

<https://code.visualstudio.com/Download>

For Windows PC, choose Windows:

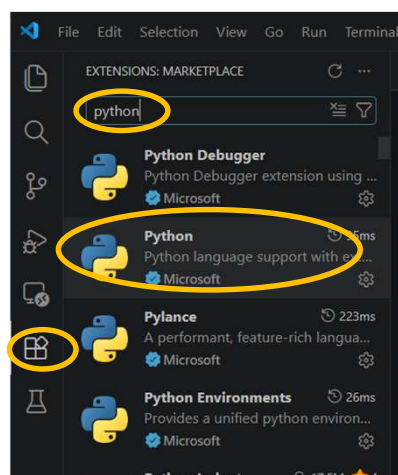


Double-click on the .exe to install the application:

 VSCodeUserSetup-x64-1.117.0.exe

7/ Install Python extension for VSC

After installing VSC, open the application. you can then install the Python extension. Go to the Extensions icon on the left panel and search for Python. Select Python, in order to install the Python extension. The Python extension is supported by Microsoft.



Set-up

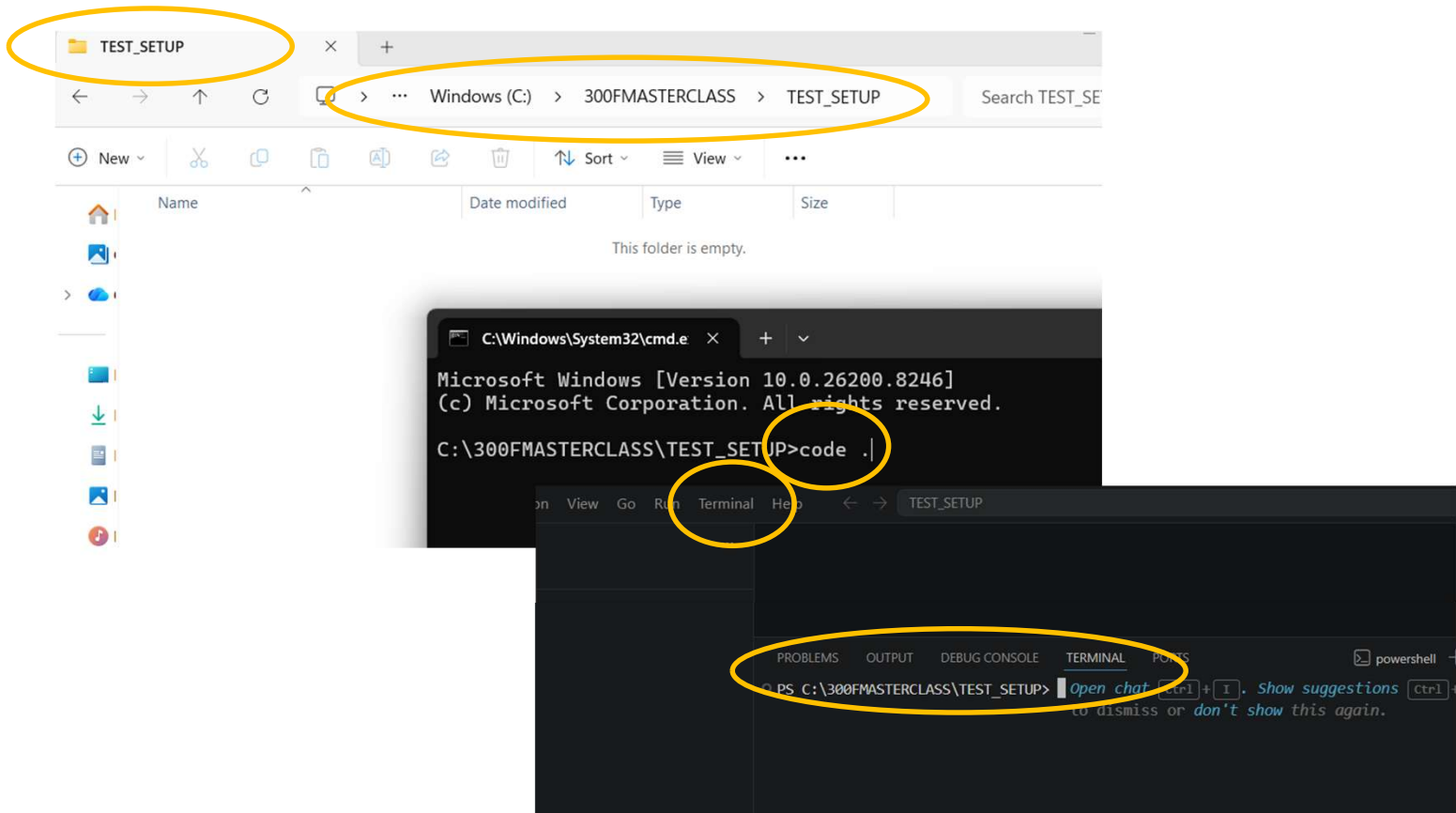
8/ Test your Python and VSC set-up -> your first python script (1/2)

Create a folder in Windows. Go inside the folder and type CMD in the search bar of the Windows explorer. Then in the black CMD box that appears, type `code .` (code space dot). This will open VSC in the location of the Windows folder.

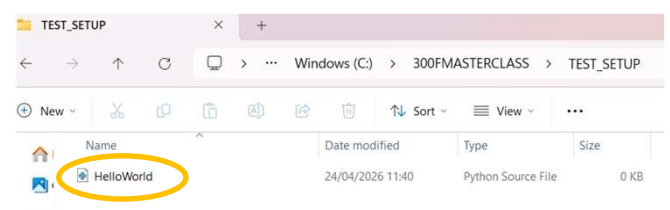
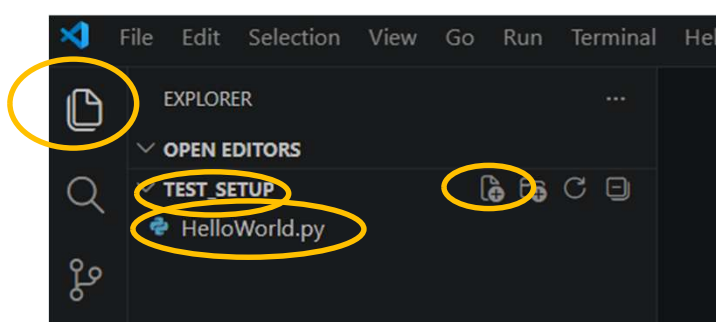
Close the Welcome and Copilot windows that pop-up inside VSC.

Go to the menu View-> Terminal. In the Terminal in VSC, you will notice that you see the same path as where you opened the CMD window.

You will also notice that the path is pre-ceded by PS. This means PowerShell. This is because the terminal is a PowerShell window, that you are viewing from inside VSC.



On the left panel of VSC, you will see that there is an **Explorer** button. Click on the Explorer button to open the left-hand panel. Here, you are looking at the Windows folders from within VSC. Hover over your Windows folder name and click on the button **New file**. Rename your file "Hello world.py". You will notice that your file appears also in your Windows folder.



Set-up

8/ Test your Python and VSC set-up -> your first python script (2/2)

Click on your HelloWorld.py file from within VSC. It will open on the right.

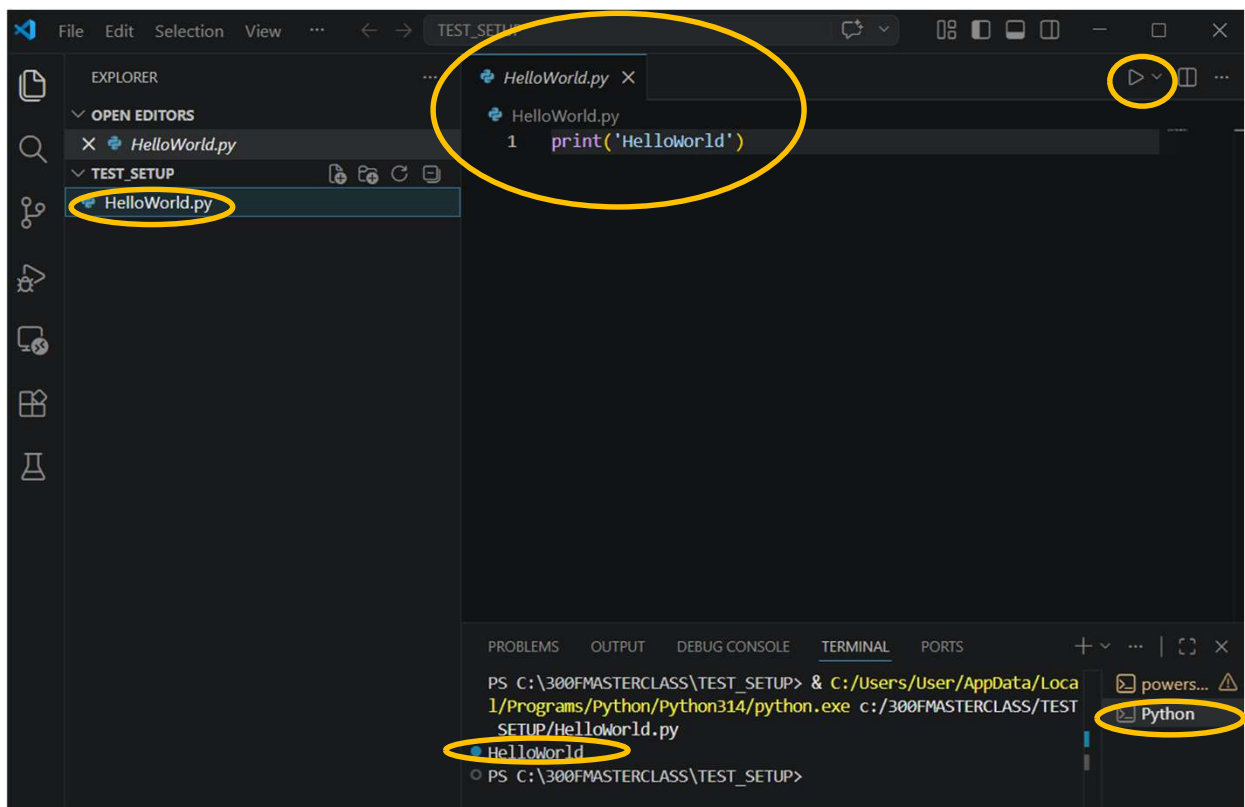
Inside the file, type `print('HelloWorld')`

Save by Choosing File->Save, or use the short-cut CTRL+Save, **or choose File and ensure that you check AutoSave (preferred).**

With HelloWorld.py still open, click on the **Run Python File** button on the top right.

Two things should happen:

- HelloWorld should be printed to the terminal.
- A Python process will appear.
 - The python process opened to run your python script. However, you are still in a PowerShell window in the Terminal, as you can see by the letters PS for PowerShell at the beginning of the line.



Well done!

You have now completed all of the steps required to use python on your machine.

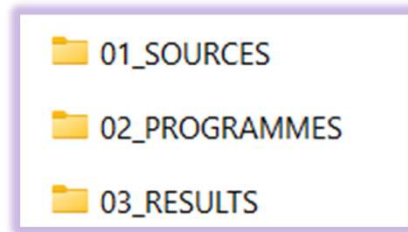
You will discover that python is much more powerful than any other data analytics tool, in terms of flexibility and reduction in number of lines of code to tackle complex tasks.

Project organization

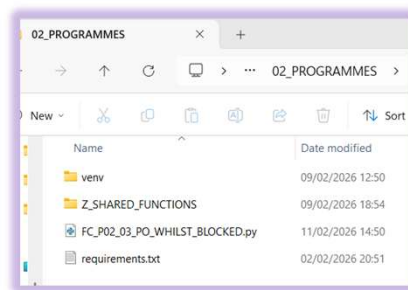
It is highly recommended to ensure that your whole team organizes their work in the same way. If anyone is sick or if you come back to a project six months later, you will be able to pick it up and understand it quickly.

Project folder organization

When creating projects, always use the following project folder structure, to avoid confusing source files with results files or program files. This ensures that you maintain a clear audit trail for your work.

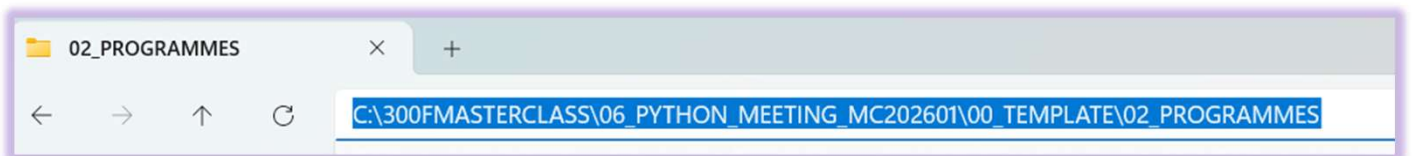


Inside 02_PROGRAMMES, you should have Z_SHARED_FUNCTIONS, where scripts that are used by other scripts should go.

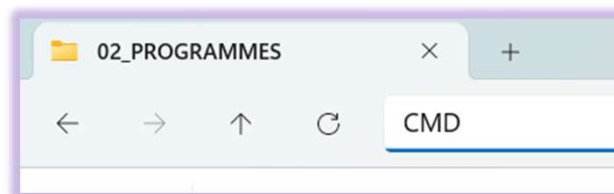


Open VSC in the correct location

In Windows, navigate to inside the 02_PROGRAMMES folder. Click in the Windows address bar.



Then type CMD to open a command window in that location.

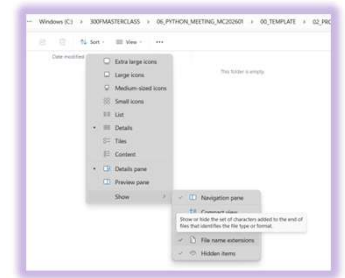


Visual Studio Code will open in the correct location for your project.

Create a toml file for your project

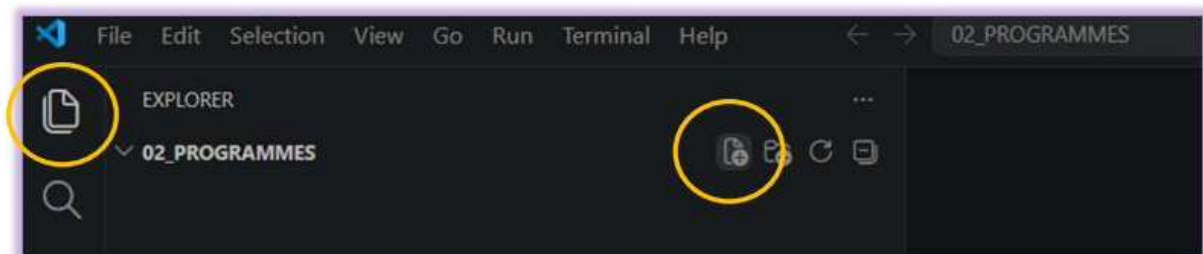
Show file extensions

When using Python, it is useful to know the types of the files that we are analysing. Ensure that you can see the file extensions in Windows. In any Windows folder, go to View -> Show and ensure that File name and extensions is ticked.



Create a .toml file

In VSC, click on the Explorer icon on the top left and then click on new file. Create a file called: pyproject.toml.



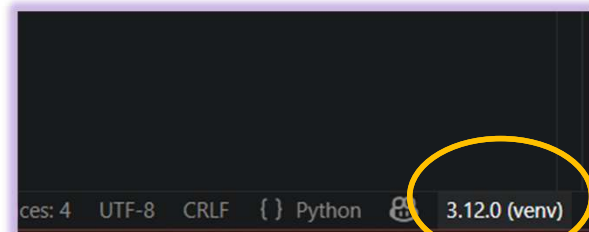
Open the file by clicking on it and then enter the following:

```
[project]
requires-python = ">=3.13,<3.15"
```

[project] indicates the project section of the .toml file. Requires-python is a variable name.

The above example is for if the Python version that you will put in your virtual environment is 3.10 (see next page). You may change the version numbers in the above example, if you intend to put a different Python version in your virtual environment.

To see the version of the python.exe that is found in your virtual environment, after creating and activating the virtual environment, as mentioned on the next page, you can type `python --version` in the terminal, or you can open a script and check the version number in the bottom right of the screen.



The .toml file is for information purposes only or for use by external systems. However, it is good practice to put the Python version in the .toml file so that anyone using your project later – on a different machine - knows which version to use when they create the virtual environment.

For example, if you create your project in version 3.10, but later, someone tries to run it with version 3.20, it is possible that the program will bug due to some incompatibilities between the Python libraries in 3.10 and the other libraries that you will import (as documented in requirements.txt – see further on), or even incompatibilities due to plain or built-in Python functions, classes or methods that are no longer found in later Python versions.

TOML stands for Tom's Obvious Minimal Language. Tom Preston-Werner (one of the co-founders of GitHub in 2008).

Create a virtual environment for your project

Open a PowerShell terminal

With VSC open in 02_PROGRAMMES, go to the menu View and then Terminal.

A terminal will appear at the bottom of the screen. Check that the prompt-line in the terminal window starts with PS.



If it starts with PS, it means that it is a PowerShell terminal. If it is not a PowerShell terminal, click on the arrow next to the + at the top-right of the terminal. Choose PowerShell.

Create and activate virtual environment

All Python projects should have their own virtual environment. We will then put all Python libraries, used in our project, inside this virtual environment. This ensures that they can always run, even if the version of Python libraries changes.

In the VSC PowerShell terminal, type the following to create a virtual environment in the folder venv: (note: here we are assuming using the Python version 3.14, but you might want to use a later version as mentioned in your .toml file (see previous page).

```
py -3.14 -m venv venv
```

The above means:

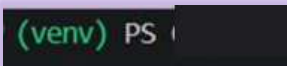
- Python: use the Python launcher
- -3.14: choose version 3.14
- -m venv: run the virtual environment module
- .venv: create the folder venv in the current folder: current folder indicated by .

The - means “this is an instruction for the command”. Be careful not to put the long dash rather than the shorter -.

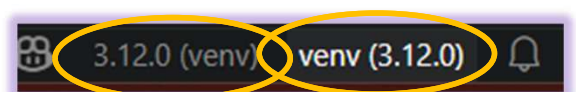
Once the venv folder has appeared in the 02_PROGRAMMES folder, we can activate it by running the activate.ps1 file. (Note the extension .ps1, because we are in a PowerShell terminal).

```
.\venv\Scripts\Activate.ps1
```

If you have successfully activated your virtual environment, you will see a green (venv) text in front of the prompt line:



You should also click/hover on the Python version number at the bottom-right of VSC: both venv() should indicate the venv in your 02_PROGRAMMES folder. If not: click and browse to correct it.



If your venv is correctly activated, your script should run using the libraries in your venv. However, if when you run your script it keeps using the wrong location for running python, you can force VSC to use the correct location by typing in the terminal, rather than using the execute button (on top right), for example:

```
.\venv\Scripts\python.exe P01_10_SUPPLIERS_IN_OFAC.py
```

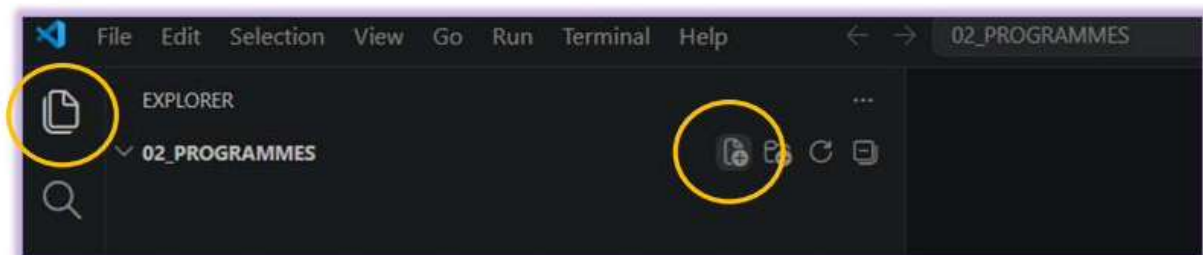
Create a requirements.txt file

Create requirements.txt

All Python projects should have a requirements.txt file. This file will hold the versions of the Python libraries that we use in our project. If we come back to our project in six months time, it may not work with different library versions. This is why we need to ensure that we record the library versions in the requirements file.

Inside VSC, click on the Explorer button on the top-left of the screen, to open the Explorer side panel on the left.

At the 02_PROGRAMMES folder level, click on the button to add a file:

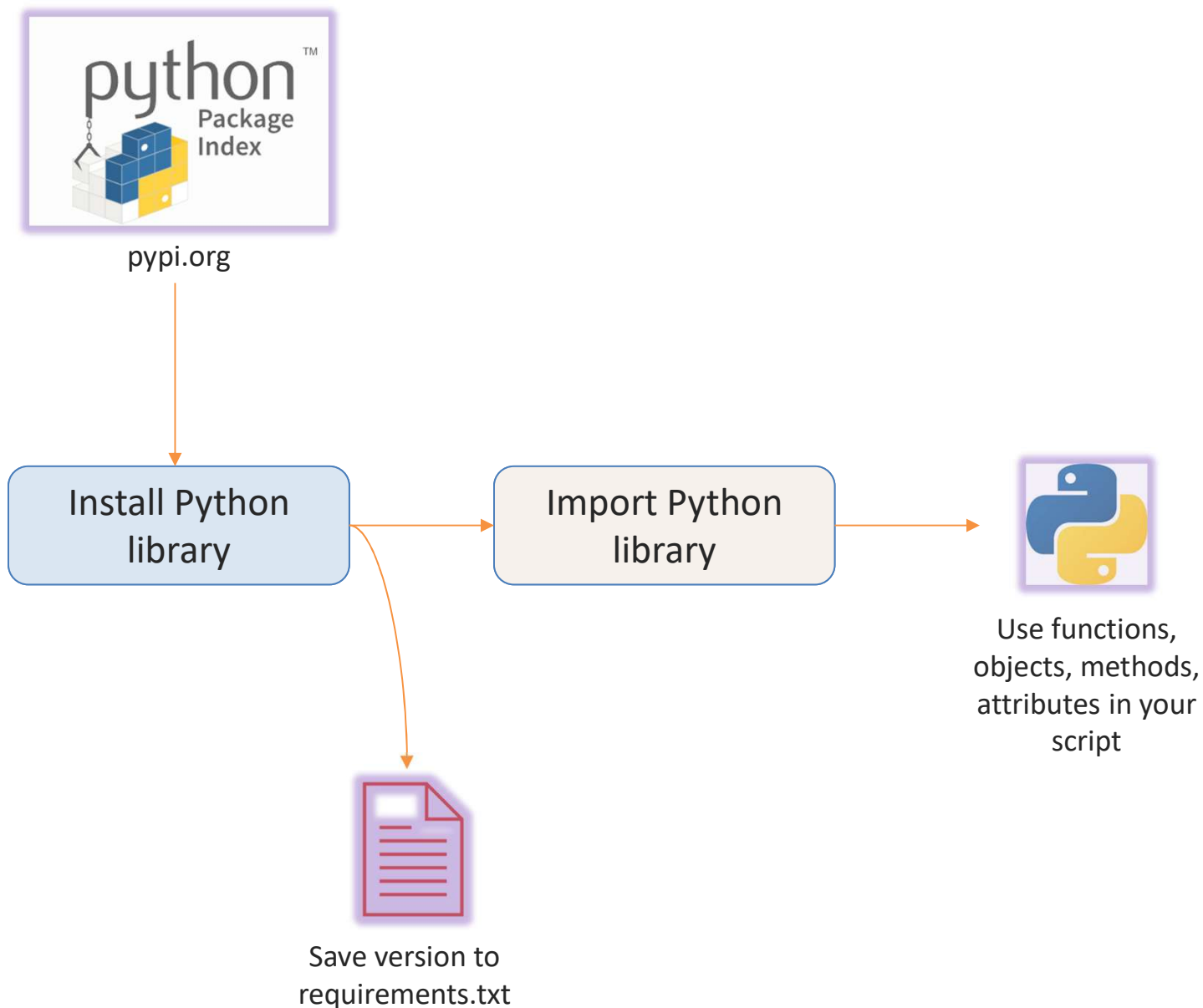


Name the file requirements.txt. As you add Python libraries to your project, you will add the mention of these libraries and their versions to the requirements.txt file, as mentioned on the next two pages.

Install and import a library

The main advantage of Python, apart from the fact that it is free and that it is the main language of Artificial Intelligence, is that it comes with endless libraries, that can enable you to do pretty much anything, with a minimal amount of code.

To benefit from all the functionality offered by the 100s 1000s libraries, we need to understand how to install, import and use Python libraries.



See the end of this document for examples of Python libraries that are used in the 300Framework.

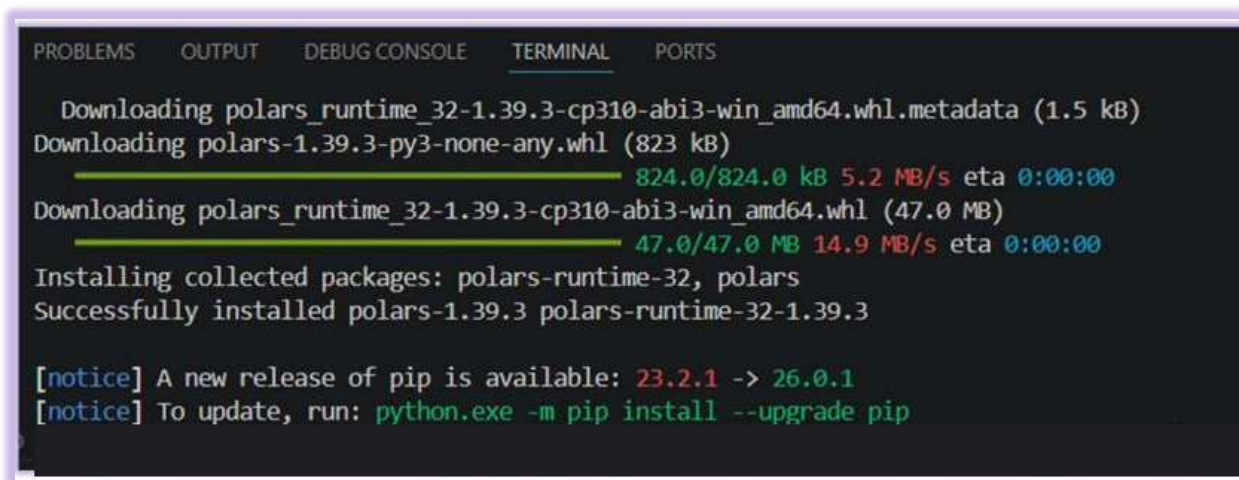
Install and import a Python library

Install a Python library

We will use Polars as an example installation of a Python library. The Python Polars library enables us to run analysis on data sets very fast. For example, filter, sort, group_by(), join, etc. We can find all of the documentation on this library here: <https://docs.pola.rs/>

In the Visual Studio Code PowerShell terminal, that is activated for your virtual environment, type the following code to download the latest version of the Polars library into your virtual environment:

```
pip install polars
```



```
PROBLEMS  OUTPUT  DEBUG CONSOLE  TERMINAL  PORTS

Downloading polars_runtime_32-1.39.3-cp310-abi3-win_amd64.whl.metadata (1.5 kB)
Downloading polars-1.39.3-py3-none-any.whl (823 kB)
  ━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━ 824.0/824.0 kB 5.2 MB/s eta 0:00:00
Downloading polars_runtime_32-1.39.3-cp310-abi3-win_amd64.whl (47.0 MB)
  ━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━━ 47.0/47.0 MB 14.9 MB/s eta 0:00:00
Installing collected packages: polars-runtime-32, polars
Successfully installed polars-1.39.3 polars-runtime-32-1.39.3

[notice] A new release of pip is available: 23.2.1 -> 26.0.1
[notice] To update, run: python.exe -m pip install --upgrade pip
```

The Polars library will be installed to the virtual environment and you can find it in the venv folder.

Record in requirements.txt for later use

Whenever we install a library, we will record it in requirements.txt. To find the version of the Polars library that you need to record, type the following in the VSC PowerShell terminal:

```
pip show polars
```

Find where it says version: ... then record the version in requirements.txt as follows:
(<libraryName>==<versionNumber>)

```
polars==1.35.0
```

To re-install libraries to your venv, in case you set-up your project on a different computer, type:

```
pip install -r requirements.txt
```

Import a library

To use a Python library in your script, you need to import it. For example:

```
import polars
```

Hide secrets

Hide secrets in an environment file and load during run-time

We often want to be able to do APIs (Application Program Interface) to different systems. Typically, when doing an API, we will need to pass a secret key. It is important that we do not put the secret key in our Python code file. Instead, we should put the key in an environment file.

An environment file should always be called .env. A .env file is typically ignored when we put our code into GitHub or other places. Furthermore, .env files are typically hidden files that you cannot see in Windows, unless you specifically show hidden files.

The .env will contain information such as:

```
ZV_ST_CHATGPT_KEY=<your-secret-key-goes-here>
```

Quotation marks are not required. There should be no space between the parameter name, the = and the value: <parameterName>=<value>.

We use the library module load_dotenv from the library dotenv to load the variables from the .env file into the process environment (the temporary memory that is used whilst the program is running, also known as the RAM).

We use the parameter override=True to mean: if the variable already exists, override it.

Once the variable is in memory, we use the getenv function of the os library to access it.

We name the variable with the same name in our script, for ease of script maintenance.

```
import os as PI_OS
from dotenv import load_dotenv as PI_LOAD_DOTENV

PI_LOAD_DOTENV(override=True)

ZV_ST_CHATGPT_KEY = PI_OS.getenv('ZV_ST_CHATGPT_KEY')
```

In this way, our secret key remains in our .env file but is not shown in our script.

if you are using GitHub, then you should always have a gitignore file in the project that mentions .env as a file that should not be sent to GitHub.

Use external variables – don't hard-code

Import and use variables from an external text file

Variables may be collected from users as inputs to a front-end dashboard. Variables may also be documented in a text file. For example, the date that the data was downloaded from the SAP system.

We put variables that are not entered by the users in the front-end dashboard, in a text file called AM_VARIABLES.txt. This text file should be located in the 01_SOURCES folder.

In the AM_VARIABLES.txt file, all variables should be written without quotation marks, except for lists of text, for example, the variable Language:

```
AnalysisStartDate=20220101
AnalysisEndDate=20231231
Language='E, EN'
MaximumRAM=20
MaximumCPU=5
```

The following Python script is FC_CONFIG_LOADER.py. It puts the values from the variables script into a dictionary.

```
from pathlib import Path as PI_PATH
from dotenv import load_dotenv as PI_LOAD_DOTENV
from dotenv import dotenv_values as PI_DOTENV_VALUES
import os as PI_OS

PI_LOAD_DOTENV(override=True)

ZV_ST_VARIABLES_FILE = (
    PI_PATH(PI_OS.getenv('ZV_ST_PROJECT_ROOT'))
    / '01_SOURCES'
    / 'AM_VARIABLES.txt'
)
ZV_DI_VARIABLES = PI_DOTENV_VALUES(ZV_ST_VARIABLES_FILE)
```

We put the FC_CONFIG_LOADER.py into the Z_SHARED_FUNCTIONS folder, and then we call it in order to create the dictionary of variables in local memory, so that it can be used by our script:

```
from Z_SHARED_FUNCTIONS.FC_CONFIG_LOADER import ZV_DI_VARIABLES
```

In the above snippet, instead of importing a function, we have imported an object.

Now, we can access the values in the dictionary. For example, we might want to get the analysis start date as a string or the maximum RAM as an integer:

```
AnalysisStartDate = ZV_DI_VARIABLES.get('AnalysisStartDate')
MaximumRam = int(ZV_DI_VARIABLES.get('MaximumRAM'))
```

(Optional) Working with GitHub

You can work with GitHub to manage and share your work within your team. GitHub will track all version changes for you and facilitate centralized storage of your projects.

When working with GitHub, we need to take time to understand the process properly so that we can effectively safeguard our team from accidentally sharing confidential information to GitHub.

1/ Create an account in GitHub

In case you don't have account, click here to create an account: <https://GitHub.com/>

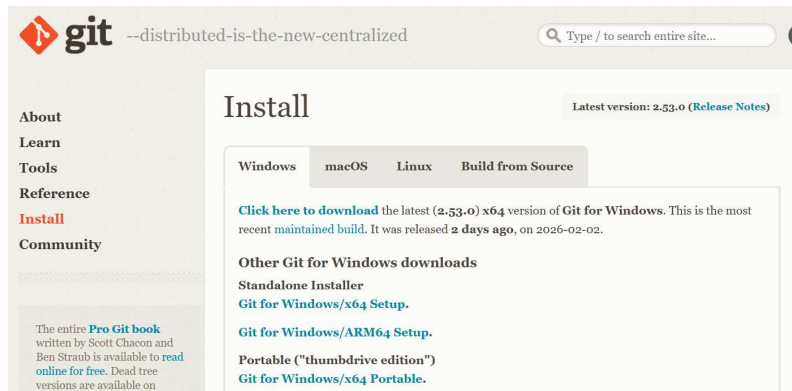
Owners of GitHub repositories will then be able to add you as members.

Note: You will need to do this step before being able to download repositories to your machine.

2/ Download and install the git application

If you would like to be able to clone projects from GitHub, it is necessary to install git on the machine where you use VSC.

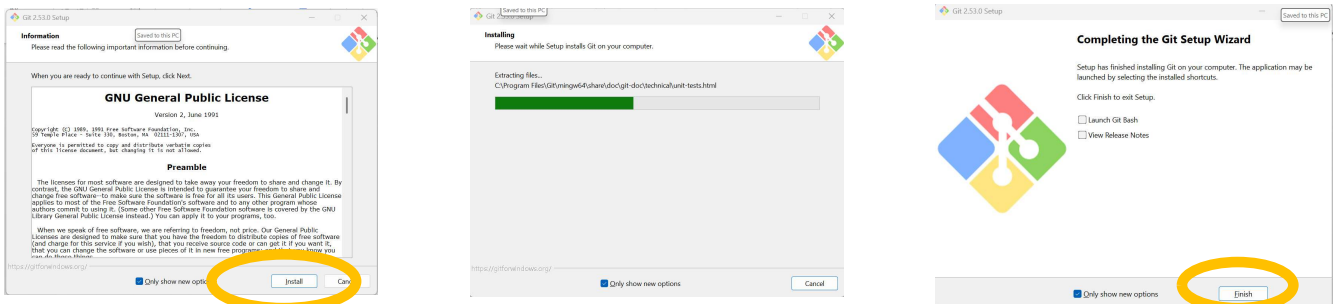
Download the git application from this link and install it on your machine: <https://git-scm.com/download/win>



The following installer will be downloaded to your machine: (Double-click on the .exe file to launch it.)

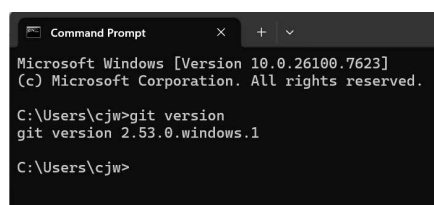


Click Install, choose Next for all questions, then at the end, un-tick both boxes and click Finish.



To test that you have correctly installed the git application, type CMD in the Windows Search bar, and then in the black box, type:

```
git version
```



If the git application is correctly installed, you will see the version number.

3/ Configure git application

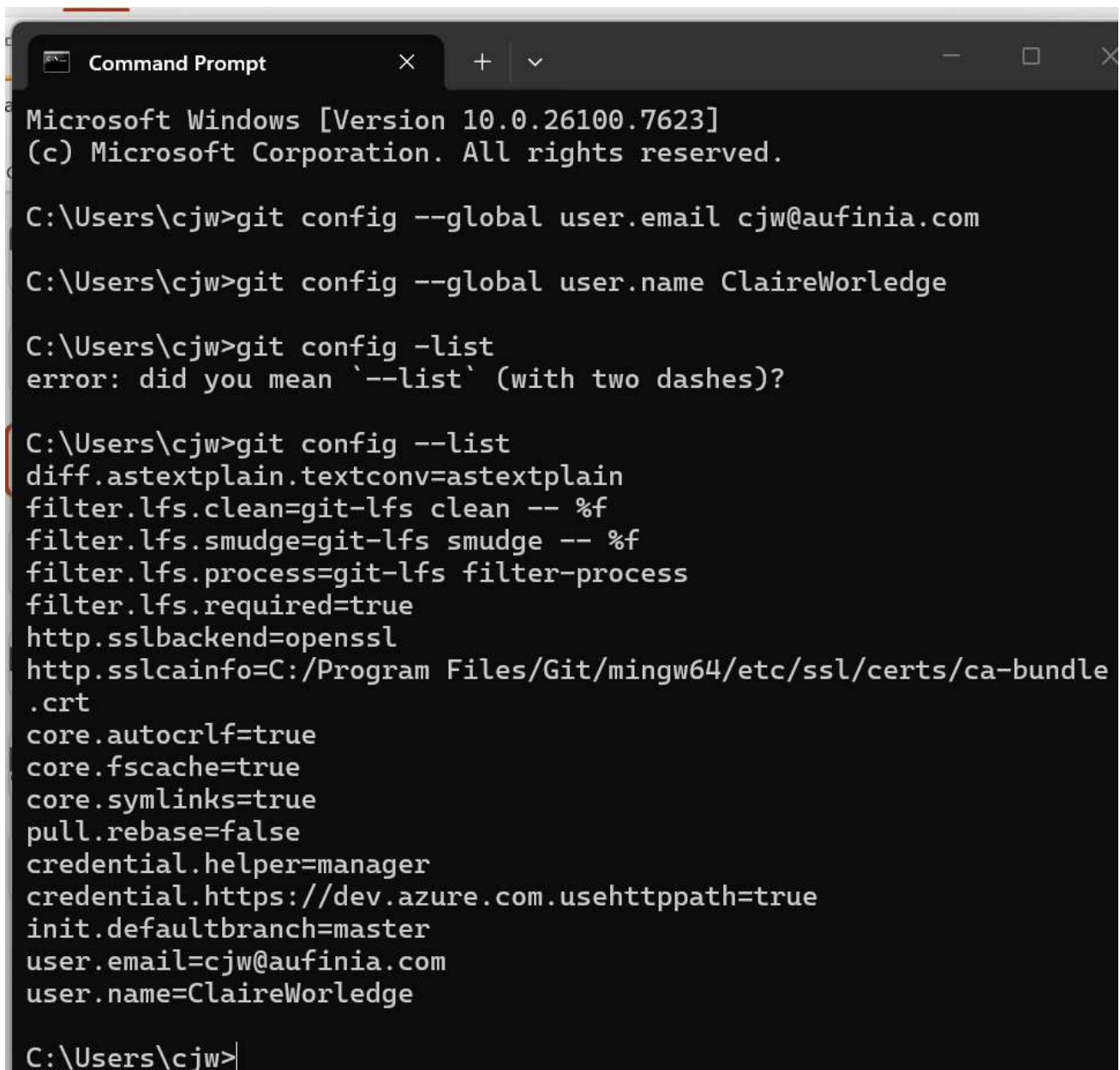
Next, we need to configure our git application, so that it can login to the GitHub account that you created in step 1.

- Type CMD in the search bar of Windows
- Then, in the black box type the following, replacing <your email> and <your name> with the email and username for your GitHub account:

```
git config --global user.email <your email>
git config --global user.name <your name>
```

- To test that your email and username were configured correctly, type:

```
git config --list
```



```
Microsoft Windows [Version 10.0.26100.7623]
(c) Microsoft Corporation. All rights reserved.

C:\Users\cjw>git config --global user.email cjl@aufinia.com

C:\Users\cjw>git config --global user.name ClaireWorledge

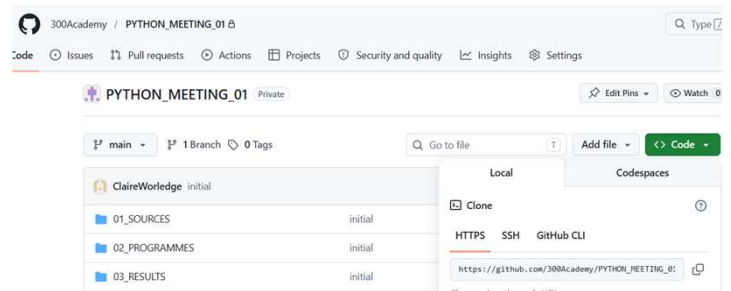
C:\Users\cjw>git config -list
error: did you mean '--list' (with two dashes)?

C:\Users\cjw>git config --list
diff.astextplain.textconv=astextplain
filter.lfs.clean=git-lfs clean -- %f
filter.lfs.smudge=git-lfs smudge -- %f
filter.lfs.process=git-lfs filter-process
filter.lfs.required=true
http.sslbackend=openssl
http.sslcainfo=C:/Program Files/Git/mingw64/etc/ssl/certs/ca-bundle
.crt
core.autocrlf=true
core.fscache=true
core.symlinks=true
pull.rebase=false
credential.helper=manager
credential.https://dev.azure.com.usehttppath=true
init.defaultbranch=master
user.email=cjl@aufinia.com
user.name=ClaireWorledge

C:\Users\cjw>
```

4/ Clone a project from GitHub (1/3)

- To clone a GitHub repository, we need to get the repository link. We can get the repository link by going to the GitHub repository and clicking on the button code. You can then copy the link.



- Go to the Windows folder, where you would like to put the Python project that you will clone.
- In the Windows search bar of that folder, type CMD
- In the black box that then appears, type the following:

```
git clone <repository link>
```

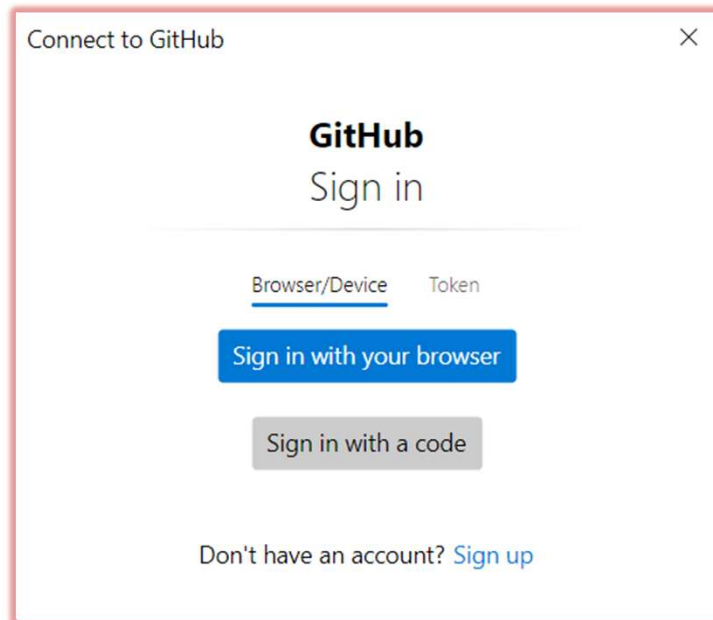
- You will see this message:

- At the same time, you will see a pop-up, as below. Click on the pop-up icon to sign in to GitHub: (Note, if you are stuck at this point, it might be because you are signed in to a different GitHub account on your machine – see the trouble shooting section).

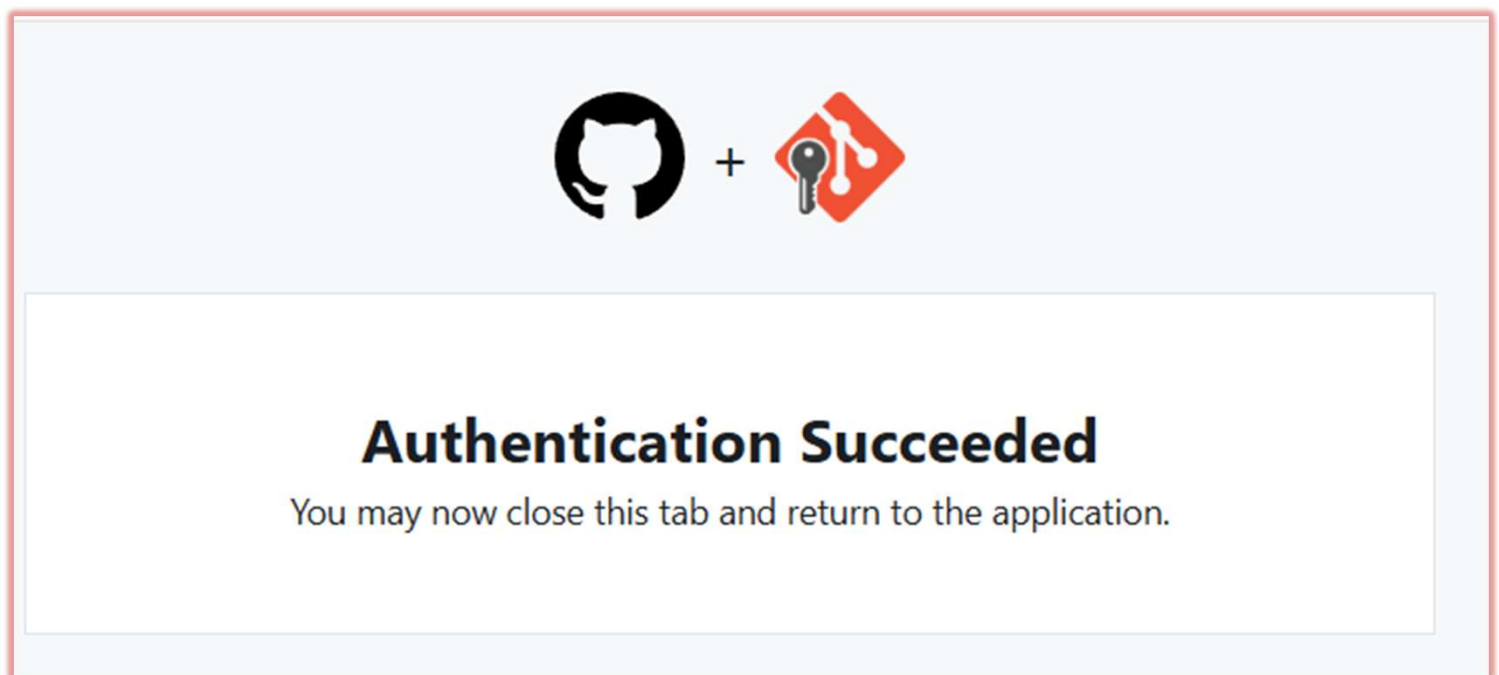


4/ Clone a project from GitHub (2/3)

- When you click on the little icon from the previous page, you will see this Window. Here you can sign in to GitHub.



- If the sign in is successful, then you will see this window:



4/ Clone a project from GitHub (3/3)

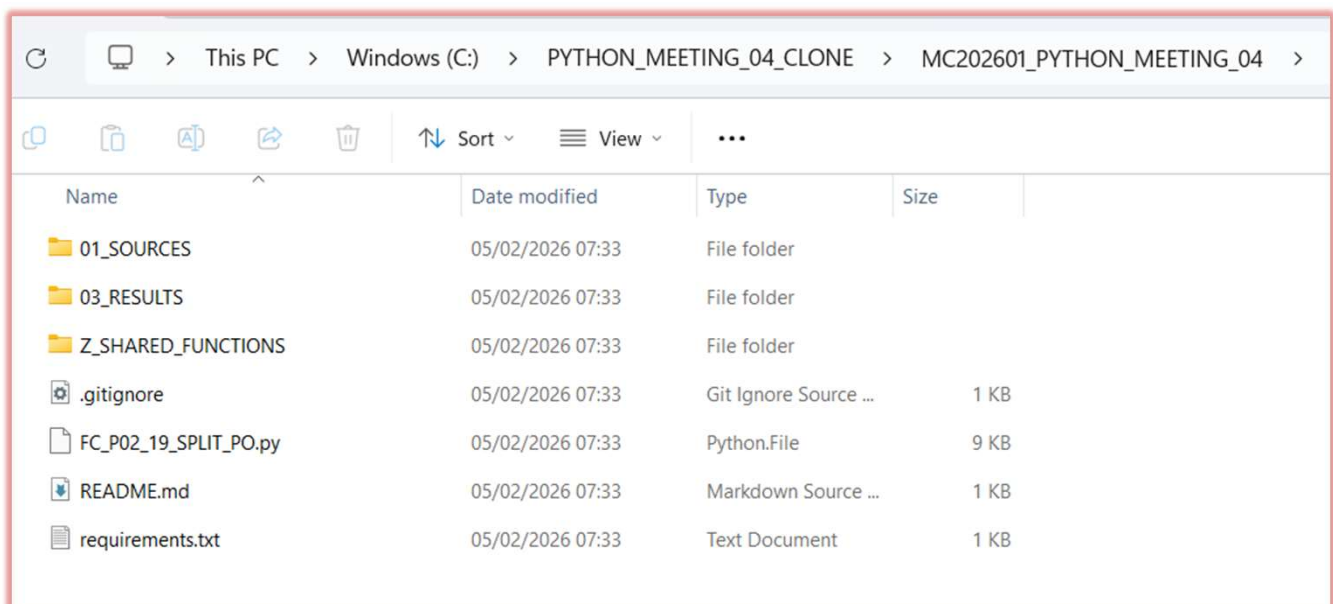
- Once you have signed-in successfully to GitHub, the folder will be downloaded to the location that you chose before:

```
C:\Windows\System32\cmd.e x + v
Microsoft Windows [Version 10.0.26100.7171]
(c) Microsoft Corporation. All rights reserved.

C:\PYTHON_MEETING_04_CLONE>git clone https://github.com/300Academy/MC202601_PYTHON_MEETING_04.git
Cloning into 'MC202601_PYTHON_MEETING_04'...
info: please complete authentication in your browser...
remote: Enumerating objects: 24, done.
remote: Counting objects: 100% (24/24), done.
remote: Compressing objects: 100% (18/18), done.
remote: Total 24 (delta 4), reused 22 (delta 4), pack-reused 0 (from 0)
Receiving objects: 100% (24/24), 393.74 KiB | 3.71 MiB/s, done.
Resolving deltas: 100% (4/4), done.

C:\PYTHON_MEETING_04_CLONE>
```

- You should then be able to see the folder structure and content in Windows:



5/ Run the code that you downloaded

When the project is downloaded from GitHub:

- Go to the 02_PROGRAMMES folder in Windows
- Type CMD in the Windows search bar and type:

```
code .
```

- Check the .toml to get the Python version.
- Create the virtual environment with the correct Python version:

```
python -<python version in .toml file> -m venv .venv
```

- Activate the venv:

```
.\venv\Scripts\Activate.ps1
```

- Install libraries:

```
pip install -r requirements.txt
```

- Update any variables that are in the AM_VARIABLES.txt file:

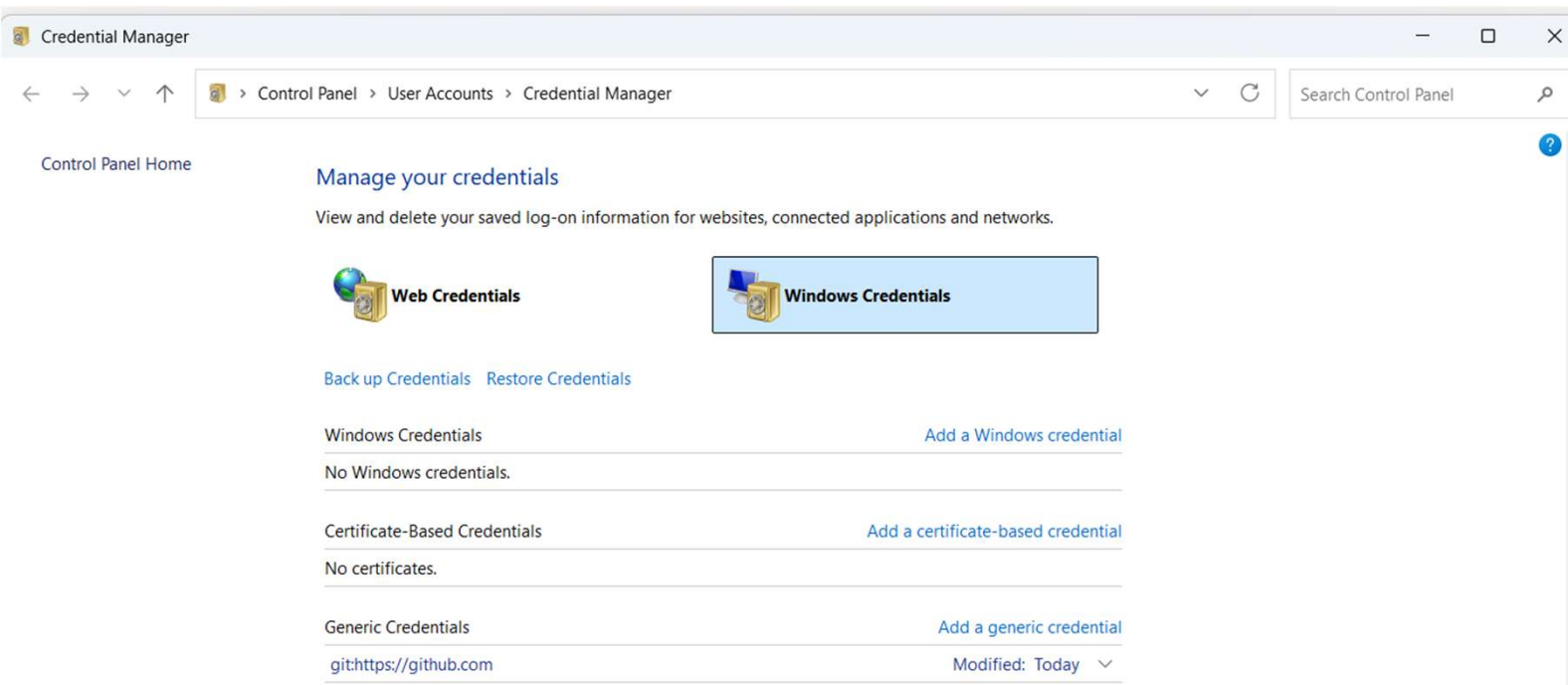
```
AnalysisStartDate=20220101  
AnalysisEndDate=20231231  
Language='E, EN'  
MaximumRAM=20  
MaximumCPU=5
```

- Run the program: we can now run our Python script, by opening the main script in the project and clicking on the play button at the top-right of VSC:



6/ Troubleshooting cloning from GitHub

- If you are not able to clone the GitHub library, it could be that you already have a GitHub account configured for your machine. In this case, you may need to remove that GitHub account (but check with your IT department first).
- To remove the GitHub account, you can type Credential Manager.
- You will then see the Credential Manager box pop-up:

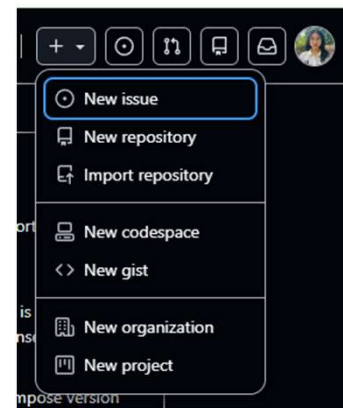


- You could try removing the GitHub account from Web credentials and starting again the process from step 1.

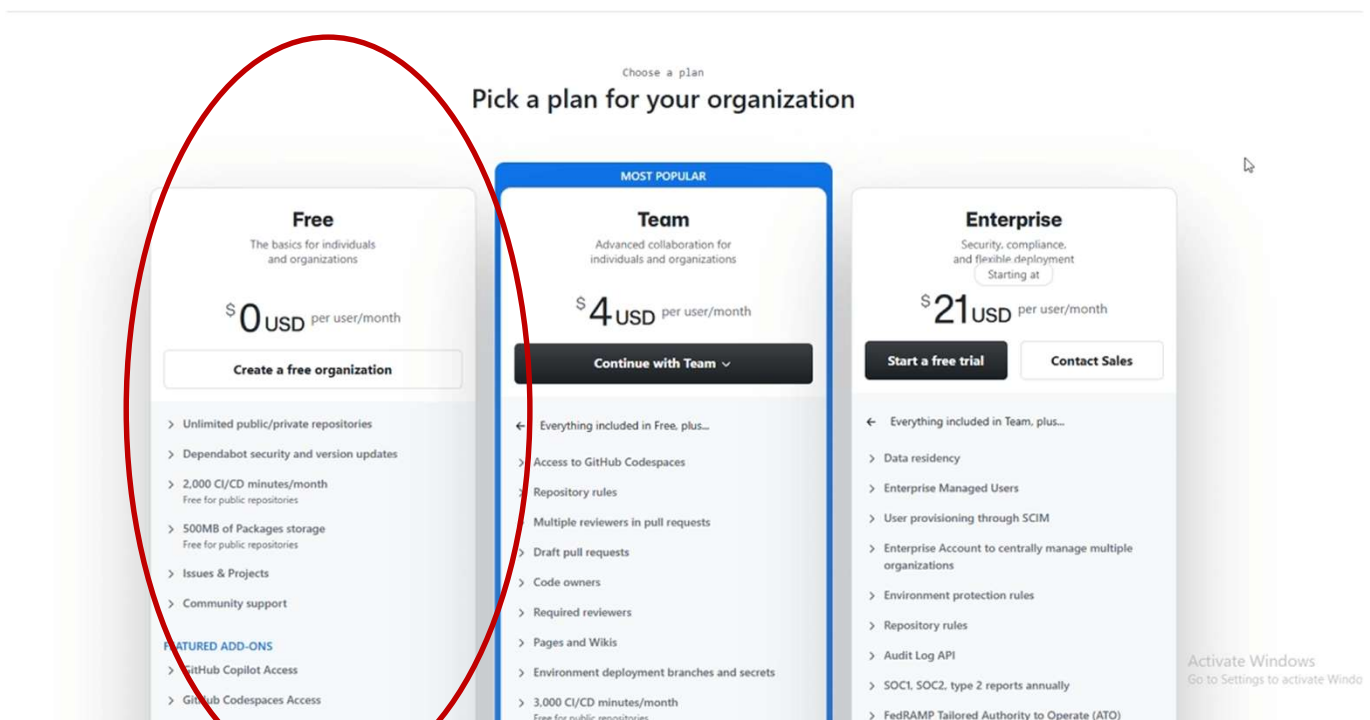
7/ Create your own organization in GitHub

- Note: never store any keys or passwords in code that is uploaded to GitHub. If you store keys, such as ChatGPT keys or the like, then it is highly likely that your code will be scrapped (even if it is private) and your keys compromised. This situation can cause you unexpected bills for the use of services, such as ChatGPT, Azure or AWS, etc.
- For this reason, it is important that only authorized team members be allowed to upload code to your GitHub.
- If you want to be able to store your projects in your own GitHub, or share them with your colleagues, then you could create your own organization

- Go to <https://GitHub.com/> to login
- Click on button '+' → choose New organization.



- Choose a plan for your organization: (you can probably use the free plan for small projects):



8/ Add name and team members to your organization

- Add information about your organization

Tell us about your organization

Set up your organization

Organization name *

This will be the name of your account on GitHub.
Your URL will be: <https://github.com/testAufinia>.

Contact email *

This organization belongs to:

☐ **My personal account**
I.e., Mai74D

☒ **A business or institution**
For example: GitHub, Inc., Example Institute, American Red Cross

Name of business or institution this organization belongs to *

This business or institution — not Mai74D (your personal account) — will control this organization.

Verify your account

- Add members (optional)

Start collaborating

Welcome to testAufinia

Add organization members

Organization members will be able to view repositories, organize into teams, review code, and tag other members using @mentions.
[Learn more about permissions for organizations](#) →

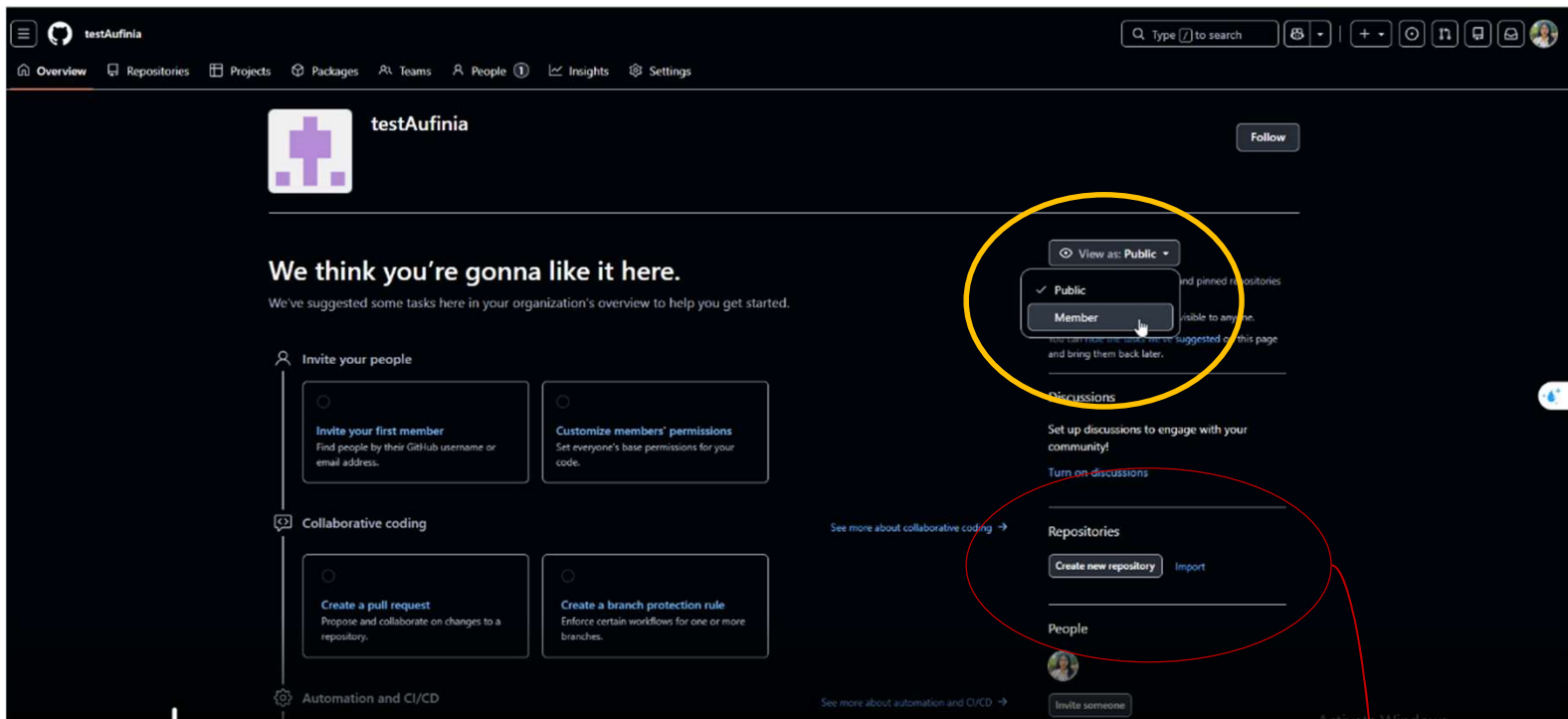
Search by username, full name or email address

Complete setup

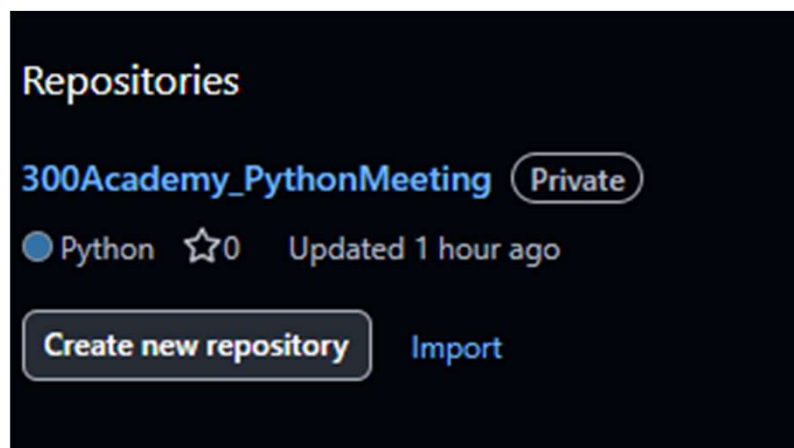
[Skip this step](#)

9/ Create a repository in your organization (1/3)

- You can view your organization as the public will see it or as members will see it. Note: this setting is only for the organization page, it does not impact the security settings of the repositories.



- To add a project (for example a Python meeting project), click on Create new repository



9/ Create a repository in your organization (2/3)

- Fill in the repository name and set to private if you only want your repository members to see this repository.

Create a new repository

Repositories contain a project's files and version history. Have a project elsewhere? [Import a repository](#).
Required fields are marked with an asterisk (*).

1

General

Owner *

 testAufinia

Repository name *

300Academy_PythonMeetingx

✓ 300Academy_PythonMeetingx is available.

Great repository names are short and memorable. How about [verbose-bassoon](#)?

Description

Contains programs related to the Python course "SAP Data Analytics Master Class Jan – Apr 2026"

95 / 350 characters

2

Configuration

Choose visibility *

Choose who can see and commit to this repository

 Private

Add README

READMEs can be used as longer descriptions. [About READMEs](#)

Off ☐

Add .gitignore

.gitignore tells git which files not to track. [About ignoring files](#)

No .gitignore

Add license

Licenses explain how others can use your code. [About licenses](#)

No license

Create repository

9/ Create a repository in your organization (3/3)

- Once you have finished creating the repository you will see the below screen.
- Click to copy the link for the next step.

The screenshot shows the GitHub interface for a newly created repository named '300Academy_PythonMeetingx' (marked as Private). At the top, there are buttons for 'Edit Pins', 'Watch' (0), 'Fork' (0), and 'Star' (0). Below the repository name, there are two main sections: 'Set up GitHub Copilot' and 'Give access to the people you work with'. The 'Quick setup' section is highlighted with a yellow circle around the copy button. It provides instructions for setting up the repository on a desktop or via command line. The '...or create a new repository on the command line' section shows a series of git commands to initialize a new repository and push it to GitHub. The '...or push an existing repository from the command line' section shows the commands to add a remote origin and push an existing repository.

300Academy_PythonMeetingx (Private)

Edit Pins Watch 0 Fork 0 Star 0

Set up GitHub Copilot
Use GitHub's AI pair programmer to autocomplete suggestions as you code.
Get started with GitHub Copilot

Give access to the people you work with
Ensure the right people and teams have access to this repository.
Manage access

Quick setup — if you've done this kind of thing before

Set up in Desktop or HTTPS SSH https://github.com/testAufinia/300Academy_PythonMeetingx.git

Get started by [creating a new file](#) or [uploading an existing file](#). We recommend every repository include a [README](#), [LICENSE](#), and [.gitignore](#).

...or create a new repository on the command line

```
echo "# 300Academy_PythonMeetingx" >> README.md
git init
git add README.md
git commit -m "first commit"
git branch -M main
git remote add origin https://github.com/testAufinia/300Academy_PythonMeetingx.git
git push -u origin main
```

...or push an existing repository from the command line

```
git remote add origin https://github.com/testAufinia/300Academy_PythonMeetingx.git
git branch -M main
git push -u origin main
```

10/ (If your token expired) Erase a token

Before being able to push to the GitHub, you need to generate a token. If you previously generated a token, and this token has now expired, you should erase the token before creating a new one.

In the Windows address bar, at the root of your project, type CMD and type:

```
code .
```

VSC should open at the root of your project. Go to View-> Terminal and type the following:

- Find out if your token is of type manager or manager-core. If it is, you will be able to erase it in the next step.

```
git config --get credential.helper
```

- Erase the token:

```
@  
protocol=s  
host=GitHub.com  
  
"@ | git credential-manager erase
```

11/ (If not already done) Create a token

Before being able to push to the GitHub, you need to generate a token.

- Login to your GitHub account and then got to: <https://GitHub.com/settings/tokens/>
- Click generate new token
- Choose Classic
- Give your token a name in Note
- Choose the token expiration (for example 90 days)
- Tick the scope: repo
- Click Generate token
- Copy the token and remember it.. Note: the token will be used when we do the push of our code to GitHub.

12/ Clone your local project to your GitHub repository (1/2)

Assuming that you have installed the git application on your machine (as mentioned in steps 2 and 3) and that you have a valid token (as mentioned in steps 10 and 11), we can now use git to copy your project to your new GitHub repository.

In the Windows address bar, at the **root** of your project, type CMD and then:

```
code .
```

optional/ alternative If you are already in the O2_PROGRAMMES location in VSC, you can go up to the root level by typing the following in the terminal:

```
cd ..
```

Visual Studio Code should open at the root of your project. Go to View-> Terminal and type the following:

- Create a hidden git folder inside your project root folder. (**optional**: Note: if you want to see this hidden folder go to any folder in Windows and select View->Show->Hidden items.)

```
git init
```

- Tell git to ignore certain folders and files. For exemple, we definitely want to ignore .venv and .env. There may be other types of files that we want to ignore:

```
@  
*.pptx  
*.pdf  
*.log  
.env  
__pycache__/  
venv/  
*.pyc  
.vscode/  
.idea/  
dist/  
build/  
*.sav  
"@ | Set-Content .gitignore
```

- **Optional**: To review what you have in gitignore, you can type:

```
code .gitignore
```

12/ Clone your local project to your GitHub repository (2/2)

- Stage all content inside the root folder, except content that is mentioned in gitignore. Note: when using git, content is staged, before it is committed (copied to GitHub).

```
git add .
```

- **Optional:** If you made a mistake and staged the wrong files, you can un-stage:

```
git restore -staged .
```

- To add a message to the commit, type the following:

```
git commit -m "initial commit"
```

- Tell git where the staged project should go:

```
git remote add origin <Repository link that you obtained by clicking on the code button in your new repository>
```

- To create a main branch in the repositor: (Note: -M is used to force, even if the branch already exists or has a different name):

```
git branch -M main
```

- To finally send the staged project (origin) to your GitHub repository (main): (Note: the -u here adds tracking. This means that next time you can just write git push, instead of git push -M main.) (Note also: The first time you do a push after creating a new GitHub token, a pop-up window will appear, requesting authentication. Choose Token and enter the token that you created in step 10.)

```
git push -u origin main
```

13/ Update changes to your GitHub repository

If you change something in your project, follow the following steps to update the files in your GitHub repository:

- Go to the root of your project

```
cd ..
```

- or

```
cd path\to\your\project
```

- **Optional:** Check what you changed: **Note: this is really useful for seeing a summary of changes since the last commit:**

```
git status
```

- Stage everything that was changed

```
git add .
```

- **Optional:** you can specify specific files:

```
git add file1.py file2.py
```

- commit the changes

```
git commit -m "Description of changes"
```

- Send the updated files to GitHub. GitHub will put the new versions and keep the history of the old.

```
git push
```

- Optional: To see what you pushed to GitHub:

```
git show --name-only
```

```
git show
```

Naming conventions & rules

To avoid errors, ensure that Python code is easy to understand and maintain, it is important that all team members follow religiously the naming conventions and rules.

OpenAI and Claude can write code for us. However, code written by AI is often convoluted and hard to follow. We can train agents to follow our naming convention rules - and we can review the code that the agents produce to ensure that the rules are respected.

Naming conventions & rules (1/3)

When writing Python code, always use the following naming conventions, to ensure your code is quick to read and understand. The below rules also help us to avoid classic errors.

<i>Example</i>	<i>Rule</i>
<code>import polars as PI_POLARS</code>	Imported libraries: PI_<LIBRARY:NAME>: (PI: PythonLibrary Import) Use Polars wherever possible, as it runs faster than Pandas.
<code>.read_csv()</code>	Code from Python libraries in lower case
UPPER	Anything that we make in upper case
<code>FC_IMPORT()</code>	Functions that we make: prefix FC_
<code>CL_BANK_ACCOUNT()</code>	Classes that we make: prefix CL_
<code>ME_DEPOSIT()</code>	Methods within classes that we make: prefix ME_
<code>FC_IMPORT (ZVFCI_ST_SOURCE_FILE, ZVFCI_ST_FOLDER, ZVFCI_ST_DELIMITER):</code>	Input variables imported into our functions: prefix ZVFCI_ (Z: custom, V: variable, FC: function, I: Input)
<code>ZV_ST_FILE_PATH</code>	Variable that we create that is of type string: prefix ZV_ST_ (Z: custom, V: Variable, ST: string)
<code>ZV_DT_START_DATE</code>	Variable that we create that is of type date: (Z: custom, V: variable, DT: date)
<code>ZV_NU_THRESHOLD</code>	Variable that we create that is of type numeric (Z: custom, V: variable, NU: numeric)
<code>ZV_LI_LANGUAGES</code>	Variable that we create that is of type list: prefix ZV_LI: (Z: custom, V: variable, LI: list)
<code>ZV_DI_APPROVAL_LIMITS</code>	Variable that we create that is of type dictionary: prefix ZV_DI (Z: custom, V: variable, DI: dictionary)
<code>ZV_DF_JE_LINES</code>	Variable that we create that is of type dataframe (for example, when we want to work on a dataframe section in a group_by): (Z: custom, V: variable, DF: dataframe)
<code>ZV_SET_TCODES</code>	Variable that we create that is of type set (Z: custom, V: variable, SET: set)
<code>ZV_SE_BKPFTCODE_TSTCTTEXT</code>	Variable that we create that is of type series (column): prefix: ZV_SE (Z: custom, V: variable, SE: series)
<code>ZV_OB_FILE</code>	Variable that we create that is another type of object – not mentioned above: prefix ZV_OB (OB: object)

Naming conventions & rules (2/3)

Example	Rule
A_<SAPTableName>	SAP table that we imported into our program
<SAPTableName>_<SAPFieldName>	Name of SAP fields that we imported to our program. For example, BSEG_DMBTR.
ZF_<SAPTableName>_<SAPFieldName>_<Action/Description>	Name of field that we created based on a SAP field, for example, ZF_BSEG_DMBTR_USD. Note: generally speaking, if we sum a field in a group_by(), we do not change the name of the field summed.
<SAPTableName>_<SAPFieldName>_<SAPTableFieldAddedTo>	Name of field that we added to another table. For example, if we added the transaction code description field (TSTCT_TTEXT) to the general ledger header (BKPF) table, then we would write TSTCT_TTEXT_BKPF. This enables us to not confuse the TSTCT_TTEXT field that is added to other tables, such as MKPF.
ZF_<SAPTableName><SAPFieldName>_<SAPFieldName>	Name of the field that we create when we concatenate the description with the code. For example, if we concatenate the transaction code description (TSTCT_TTEXT) with the transaction code (BKPF_TCODE), then we would write ZF_BKPF_TCODE_TTEXT
ZF_ST_<Description>	For example, ZF_ST_JE_TYPE. If we create a string field to describe the journal entry type: such as manual / normal.
<ScriptNumber>_<99>_TT_DF_<SAPTable>_<Action/Description> for example: B01_01_TT_DF_BSEGBKPF_FILTER	Temporary dataframes (tables) that we make: prefix: script number that made the dataframe, _<99>_DF_TT (<99>: chronological order of creation, DF: dataframe, TT: Temporary Table)
<ScriptNumber>_<99>_IT_DF_<SAPTable>_<Description> for example: B01_12_IT_DF_BSEGBKPF_CUBE	Dataframes that we will use later: prefix: script number that made the dataframe, _<99>_DF_IT (<99>: chronological order of creation, DF: dataframe, IT: Interpreted Table)
<ScriptNumber>_<99>_RT_DF_<SAPTable>_<Description> for example: B01_13_RT_DF_BSEGBKPF_TOTALS	Dataframes that we will later present in the 300Framework dashboard and that show us information, such as totals per supplier.
<ScriptNumber>_<99>_XT_DF_<SAPTable>_<Description> for example: B01_13_XT_DF_BSAIK_DUP_PAY	Dataframes that we will later present in the 300Framework dashboard and that show us exceptions to internal control rules, such as duplicate payments.

Naming conventions & rules (3/3)

Example	Rule
<pre>B01_01_TT_BSEG_COUNT = (A_BSEG .group_by(BSEG_BUKRS, BSEG_GJAHR, BSEG_BELNR) .agg(PI_POLARS.col('BSEG_BELNR') .n_unique() .alias('ZF_BSEG_BELNR_COUNT')))</pre>	Write code vertically for fast half-screen review. We often split our screen in half to compare two sets of code, or compare code to something else. For example, to compare code from ChatGPT or Claude to our code-base. Therefore, for quick review, we write all of our code vertically. Chain with . for continued steps. Use the Polars library, wherever possible, for fast execution.
<pre># 1/ Import</pre>	All steps should have brief comments. Comments should be chronologically numbered.
<pre>Obsolete code snippets</pre>	Obsolete code must be deleted. This is to ensure that we don't waste time reviewing it and so that we do not confuse AI agents.
<pre>FC_<ActionOrServiceDescription> in Z_SHARED_FUNCTIONS - no repetition</pre>	Code must not be repeated. Repeated functions, such as import text file, export to Excel, export to SQL database, must go in functions, found in separate Python scripts. All such functions must be put in a separate Python script in Z_SHARED_FUNCTIONS. This ensures that updates can be done in one central location. For example, if you want to improve the way you import text files, you want to make sure that you only have to modify the import script one time.
<pre>Technical names under presentation level, friendly names on the presentation level.</pre>	Data that is provided to end-user presentations(dashboards, etc) must be provided with technical names. Friendly names must be shown to the user on the presentation level. (on the dashboard). This is to ensure that dashboards are easy to read, but that information relating to the technical origin of the data remains up until the last level, for quick maintenance and review purposes.
<pre>pyproject.toml</pre>	All projects must have a pyproject.toml that indicates the Python version that was used when the virtual environment was created.
<pre>.venv</pre>	All projects must have a virtual environment folder, in which the Python.exe inside the .venv/Scripts is the same as that mentioned in the .toml.
<pre>requirements.txt</pre>	All projects must have a requirements.txt file that lists all of the libraries that need to be installed for the scripts in the project to work.
<pre>.env</pre>	All projects must have a .env file for storing secrets. Keys and passwords should never be included in the Python script itself.
<pre>no hard-coding</pre>	Variables should never be hard-coded. All variables must come from a user-input or a parameters file. This is to avoid future errors due to values, such as start date, or end date, etc. that are forgotten inside a script.
<pre>Short scripts - one function per script</pre>	Scripts should be written with one purpose only and they should be as short as possible. This is to ensure that they are easy to maintain, but also to ensure that the RAM is released between each operation, to prevent RAM overload.

Python components

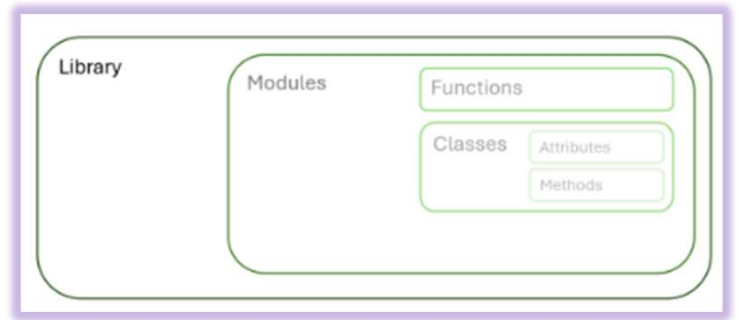
AI models such as OpenAI and Claude can produce Python code for us. However, they often make mistakes. When we have a basic grounding in Python, we can easily follow what the AI produces and notice where mistakes are made.

To be able to read Python code quickly, it helps if we can visualize the structures.

How is Python organized?

Python is organized as follows:

- Libraries
 - Modules / sub-modules (namespaces)
 - Functions
 - Classes
 - Attributes
 - Methods



- Functions: do actions.
- Classes: create objects.
- Objects: have attributes and methods
- Methods do actions on objects or attributes of objects in the same class.

We can download libraries and use their functions and classes.

We can also make our own functions and classes.

Functions

Functions do an action.

A function can ally take objects as parameters.

A function can ally return objects as results.

Example: Function that does not take or return any objects:

```
def FC_HELLO_WORLD():  
    print("hello world")
```

Example: Function that takes an object:

```
def FC_WELCOME(ZVFCI_ST_NAME):  
    print(f"Hi {ZVFCI_ST_NAME}, Welcome to the 300Framework Masterclass")
```

This function takes the object type string variable as an input.

Example: Function that takes and returns objects:

```
def FC_BALANCE(ZVFCI_DF_BSID):  
    ZV_DF_BSID_KUNNR_TOT = (  
        ZVFCI_DF_BSID  
        .group_by(  
            BSID_KUNNR  
        )  
        .agg(  
            PI_POLARS.col('ZF_BSID_DMBTR_USD_SIGNED')  
            .sum()  
        )  
    )  
    return ZV_DF_BSID_KUNNR_TOT
```

This function takes a dataframe object (customer outstanding items) in its signature and returns a dataframe object (total outstanding per customer)

Classes

Below, is an example of a custom class, that we can create, in order to understand what a class is and how to use it.

```
Class CL_BANK_ACCOUNT:
    def __init__(self, ZVCLI_ST_OWNER, ZVCLI_NU_BALANCE):
        self.ZV_ST_OWNER = ZVCLI_ST_OWNER
        self.ZV_NU_BALANCE = ZVCLI_NU_BALANCE

    def ME_DEPOSIT(self, ZVMEI_NU_AMOUNT):
        self.ZV_NU_BALANCE = self.ZV_NU_BALANCE + ZVMEI_NU_AMOUNT

    def ME_WITHDRAW(self, ZVMEI_NU_AMOUNT):
        if ZVMEI_NU_AMOUNT <= self.ZV_NU_BALANCE:
            self.ZV_NU_BALANCE = self.ZV_NU_BALANCE - ZVMEI_NU_AMOUNT
        else:
            print("not enough money")
```

Using this class, we can create a bank account object: ZV_OB_ACCOUNT. This object is an instance of the CL_BANK_ACCOUNT class.

We can use this bank account object to deposit and withdraw, using the methods of the CL_BANK_ACCOUNT class.

```
ZV_OB_ACCOUNT = CL_BANK_ACCOUNT('Claire', 1000)
ZV_OB_ACCOUNT.ME_DEPOSIT(250)
ZV_OB_ACCOUNT.ME_WITHDRAW(100)
```

Notice how we replace `self` with ZV_OB_ACCOUNT.

Objects

In Python, everything is an object. Dataframes, columns, variables, lists and dictionaries are all examples of objects. The functions that we can use on an object, depend on the object class definition.

For example, a dataframe can be a Polars dataframe or a Pandas dataframe. If the dataframe is a Polars dataframe, then we can only use Polars functions on it and not Pandas functions.

Dataframe

A dataframe in Python is almost the same as a table. A dataframe has columns (also referred to in Python literature as Series) and rows.

Column

We use `PI_POLARS.col()` to access a column in a Polars dataframe. Once accessed, functions, methods and operations can be done on all items of the column. For example, `PI_POLARS.col('ZF_ST_MY_COL').str.strip_chars()` removes all leading/ trailing spaces from every item in `ZF_ST_MY_COL`.

Variable

A variable is one value. We can think of a variable as a box with one value inside. When working on variables, we can use normal Python functions, rather than Polars functions.

List

In Python, a list is written with `[]` brackets. A list can contain any number of items. Each item can be of a different object class (dataframe, variable, list, dictionary). Items in the list can be accessed based on the index (position in the list). The first position in a list has the index 0.

A list is written as `[<object>, <object>, <object>]`, where object can be of any type (text, numeric, list, dictionary, ...). Objects can be of different types within the same list. The same object can occur more than once.

```
ZV_LI_MY_LIST = [object, object, object, object]
```

Each position in the list has an index, starting at 0:

- 1st object: index 0
- 2nd object: index 1
- 3rd object: index 2
- 4th object: index 3

Objects

Dictionary

In Python, a dictionary is written with {} brackets. A dictionary can contain any number of items. Each item can be of a different object class (dataframe, variable, list, dictionary, ...). Items in the dictionary can be accessed based on the key, which is the name of the item.

A dictionary is written as {"<keyName>": <object>}, where object can be of any type (text, numeric, list, dictionary,...). Objects can be of different types within the same dictionary. We typically use nested dictionaries.

```
ZV_DI_MY_DICTIONARY = {  
    'key': {  
        'key': {  
            'key': myObject  
        }  
    }  
}
```

Group_by()

When we use group_by(), we obtain a group_by() object. A group_by() object is a dataframe that has been divided up into sub-dataframes.

```
ZV_OB_GROUPBY = (  
    ZV_DF  
    .group_by(  
        [  
            'Material',  
            'Month',  
        ]  
    )  
)
```



Material	Month	Posting date	Input date	Input time	Document	Item	Quantity	Value
Bolts	Jan 2025	1 Jan 2025	1 Jan 2025	00:01:30	00001234	0001	10	200
Bolts	Jan 2025	2 Jan 2025	2 Jan 2025	00:05:20	00001235	0001	3	60
Bolts	Jan 2025	3 Jan 2025	3 Jan 2025	00:16:10	00001236	0001	2	400
Bolts	Jan 2025	25 Jan 2025	25 Jan 2025	00:01:30	00001237	0001	14	280
Bolts	Feb 2025	1 Feb 2025	1 Jan 2025	00:01:30	00001238	0001	5	1000
Bolts	Feb 2025	2 Feb 2025	1 Jan 2025	00:05:20	00001241	0001	6	1200
Bolts	Feb 2025	3 Feb 2025	1 Jan 2025	00:16:10	00001242	0001	17	3400
Bolts	Feb 2025	25 Feb 2025	1 Jan 2025	00:01:30	00001256	0001	2	400
Bolts	Mar 2025	1 Mar 2025	1 Mar 2025	00:01:30	00001281	0001	21	4200
Bolts	Mar 2025	28 Mar 2025	28 Mar 2025	00:05:20	00001290	0001	32	6400

As soon as we apply a method on the group_by() object, it is transformed back into a dataframe. Typically, we apply the method agg() to a group_by() object and use other methods on columns inside agg():

```
.group_by(...)  
.agg(  
    [  
        PI_POLARS.col('..').sum().alias(''),  
        PI_POLARS.col('..').mean().alias(''),  
        PI_POLARS.col('..').first().alias(''),  
        PI_POLARS.col('..').n_unique().alias('')  
    ]  
)
```

In the above code, you can see that we can use the Polars col() function to bring up a column and then we can apply chained methods on that column. This gives us a lot of flexibility.

Custom

As mentioned above – a class is used to define an object. You can make your own class and so your own objects. Your objects can have their own customized attributes and methods. An attribute is also an object: for example, a string variable that holds a name.

Datatypes

Datatypes determine the function, method, operands

The functions, methods or operands (+,-,*,&,|) that we can use on an object, depend on the object and also on the type.

If our object is a variable or a column, it can also be of type, string, date, numeric or Boolean.

If our object is a list, the items in the list can be of type string, date, numeric or Boolean.

Sometimes, we need to convert between different types. The function that we use to do the conversion depends on the class and the type.

For example, the function to convert a variable of type string to a variable of type numeric is different to the function that we use to convert a Polars column of type string to numeric or a Pandas column of type string to numeric.

Python syntax

Indent

In Python, code has to be written on the correct tab indent from the left. For example, everything after `if... :` has to be indented by one tab, if it is to be done only when the `if` condition is true.

Parenthesis ()

Parenthesis are used to allow code to go on a separate line – so that we can write code vertically – as mentioned in our naming convention. Parenthesis are also used on function names. They hold the signature, which is the list of variables that can be fed to the function.

:

The colon is used after the name/ expression of **functions, for, if**. For example, `if ... : for ... :`
`FC_MY_FUNCTION(...):`

Square brackets []

Square brackets are used to hold items in a list. Lists are typically seen when a function requires a list as an input parameter. For example, `group_by([])`, which takes a list of column names.

Curly braces {}

Curly braces are used to hold items in a dictionary. Dictionaries are typically seen when a function requires a dictionary as an input parameter. For example, `rename({})`, which takes a dictionary that has the current column names as the keys and the new column names as the values.

.

. means "apply my function or method to the object or attribute before the dot". Or "give me the attribute of an object": `<object>.<attribute>`. Or "access the function that is in the library or name-space of ...

```
ZV_DF = (  
    ZV_DF  
    .with_columns(  
        PI_POLARS.col('BSEG_HKONT')  
        .str.strip_chars()  
    )  
)
```

This means: my dataframe object is equal to my dataframe object, to which you should apply the `with_columns()` method. Use the `col()` function from the Polars library to get the column `BSEG_HKONT`, to which you should apply the method `strip_chars()`, which is part of the `str()` namespace. Note: `col()` is a function, whereas `with_columns`, `str` and `strip_chars()` are methods. Functions that are not built-in (those coming from imported libraries) have the library name in front of the dot. Methods have the object name in front of the

.

&, |, ==, !=, ~

Means: **AND, OR, is equal to, is not equal to, NOT**

=

Means assign. For example, `ZV_ST_MY_NAME = 'Claire'`

Code snippets

Understanding Python code, enables us to be able to effectively review the work of AI or team members to check that our reasoning is sound, given the business context.

Having a clear knowledge of Python helps us to have a mental toolbox, that we can reach into, in order to create interesting analytics.

It is a bit like a chess player. Everyone knows the basic rules of chess. But those that study the chess books have examples of patterns in their mind that they can use to their advantage. Successful chess players can use different game patterns to creatively win against their opponent.

These code snippets can be considered like these chess patterns. Successful team members can use different python code patterns to find the best way to do any analytics that they require.

You can also find full documentation of Python functions here:

<https://docs.python.org/3/library/functions.html>

Dataframe: print to view / access columns

Print top of a file to the terminal

We often want to print to terminal to view a dataframe during processing: We use f'{' to call code, or expose the content of variables, whilst printing:

```
print(f'ZV_DF: {ZV_DF.head(10)}')
```

The above code prints the first 10 rows of the dataframe to the terminal.

Obtain the list of columns in a dataframe

We often want to know the list of columns in a dataframe, for example to check that we have all of the columns that we need for our audit test:

```
ZV_LI_COLS = ZV_DF.columns
```

In the above code, columns is an attribute of the object dataframe. ZV_LI_COLS is a list of strings, which are the names of the columns.

The syntax <object>.<attribute> enables you to access attributes inside your object.

Because we are accessing the attribute and not calling a method, we do not need paranthesis. So we do not write ZV_DF.columns().

Dataframe: process columns

Process all columns in a dataframe

We may want to process all columns in a dataframe. For example, add the prefix of the SAP table name:

```
for ZV_ST_COL_NAME in ZV_DF.columns:
    ZV_DF = (
        ZV_DF
        .rename(
            {
                ZV_ST_COL_NAME: f'{ZVFCI_ST_PREFIX}_{ZV_ST_COL_NAME}'
            }
        )
    )
```

Process all columns in a dataframe, if found in list

We may want to process all columns in a dataframe, only if they are found in a list. For example, convert some of the columns to numeric datatype:

```
for ZV_ST_COL_NAME in ZVFCI_LI_NUM_COLS:
    ZV_DF = (
        ZV_DF
        .with_columns(
            [
                PI_POLARS.col(ZV_ST_COL_NAME)
                .cast(PI_POLARS.Float64, strict=False)
                .alias(ZV_ST_COL_NAME)
            ]
        )
    )
```

Dataframe: access item

Obtain a value from a dataframe

We often want to get a value from a dataframe. For example, we might want to get the latest date found. We can use `select().max()` and `item()`:

```
ZV_ST_BSAID_BUDAT_LATEST = (  
    ZV_DF_BSAID  
    .select(  
        PI_POLARS.col('BSAID_BUDAT')  
        .max()  
    )  
    .item()  
)
```

Dataframe: filter based on value/ variable

Filter on a hard-coded value

We often want to filter a dataframe to not include rows with blank text or zero values:

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .filter(  
        (PI_POLARS.col('EKPO_MATNR') != '') &  
        (PI_POLARS.col('EKPO_NETWR') != 0)  
    )  
)
```

Filter on a variable

We often want to filter a dataframe based on a variable. The variable could be an information that is entered into a parameters file, or to the dashboard, by an end-user.

For example, in the below code, we filter on the year of the start and end date that is passed to our program.

Note: that we use [:4]. Here : means slice. Nothing before the slice means start from the beginning. 4 means up to the 4th character.

Note: for variables we can use [:4] to get the first four digits. However, for columns (series), we need to use the Polars method from the str namespace: ZF_ST_MYCOLUMN.str.slice(0,4).

```
ZV_DF_LFC1 = (  
    ZV_DF_LFC1  
    .filter(  
        (PI_POLARS.col('LFC1_GJAHR') >= ZV_ST_START_DATE[:4]) &  
        (PI_POLARS.col('LFC1_GJAHR') <= ZV_ST_END_DATE[:4])  
    )  
)
```

Dataframe: filter based on list

Found in a list

We often want to know if our general ledger account is found in a list. For example, we might have a list of cash accounts and we want to filter our general ledger to see if it contains those accounts:

```
ZV_LI_BSEG_HKONT_CASH = ['55012345', '550123333']

ZV_DF_BSEG = (
    ZV_DF_BSEG
    .filter(
        PI_POLARS.col('BSEG_HKONT')
        .is_in(
            ZV_LI_BSEG_HKONT_CASH
        )
    )
)
```

Not found in a list

Sometimes, we want to filter on records that are not found in a list. For example, we may want to obtain all suppliers that are not intercompany or personnel. We can use ~, for not.

```
ZV_LI_KTOKK_INTERCO = ['INT', 'GROUP', 'PERS']

ZV_DF_LFA1 = (
    ZV_DF_LFA1
    .filter(
        ~(
            PI_POLARS.col('LFA1_KTOKK')
            .is_in(
                ZV_LI_KTOKK_INTERCO
            )
        )
    )
)
```

Dataframe: group and agg()

Group rows – to calculate sum, mean, etc.

We often want to do operations on groups of rows. For example, the average price per material, or the total value per material. This is also another way of making sure a key is unique before a join.

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .group_by(  
        'EKPO_MATNR'  
    )  
    .agg(  
        [  
            (  
                PI_POLARS.col('ZF_EKPO_NETPR_USD') *  
                PI_POLARS.col('EKPO_PEINH') *  
                (  
                    PI_POLARS.col('EKPO_BPUMN') /  
                    PI_POLARS.col('EKPO_BPUMZ')  
                )  
            )  
            .mean()  
            .alias('ZF_EKPO_NETPR_USD_MEAN'),  
            PI_POLARS.col('ZF_EKPO_NETWR_USD')  
            .sum()  
        ]  
    )  
)
```

In the code above, we can see that `group_by()` is a method that works on a Polars dataframe to return a grouped Polars dataframe. Whereas, `agg()` is a Polars method that works on a grouped Polars dataframe.

The `group_by` method takes a string representing a column name, or a list of strings for more than one column name.

The `agg()` method takes a Polars expression or a list of Polars expressions.

Dataframe: group and agg()

Group rows – to list out words

We often want to list out the words in a string column. For example, we often want to obtain the list of accounts on debit and credit for our journal entries:

```
ZV_DF_BSEG = (
    ZV_DF_BSEG_ACC_SUM
    .sort(
        [
            'BSEG_BUKRS',
            'BSEG_GJAHR',
            'BSEG_BELNR',
            'BSEG_SHKZG',
            'ZF_BSEG_DMBTR_SUM',
        ],
        descending=[False, False, False, False, True]    # <- force numeric desc on the sum
    )
    .group_by(
        [
            'BSEG_BUKRS',
            'BSEG_GJAHR',
            'BSEG_BELNR',
        ],
        maintain_order=True
    )
    .agg(
        [
            (
                PI_POLARS.col('BSEG_HKONT')
                .filter(PI_POLARS.col('BSEG_SHKZG') == 'S')
                .cast(PI_POLARS.Utf8)
                .implode()
                .list.slice(0, 5)
                .list.join('_')
                .alias('ZF_BSEG_HKONT_DEBIT')
            ),
            (
                PI_POLARS.col('BSEG_HKONT')
                .filter(PI_POLARS.col('BSEG_SHKZG') == 'H')
                .cast(PI_POLARS.Utf8)
                .implode()
                .list.slice(0, 5)
                .list.join('_')
                .alias('ZF_BSEG_HKONT_CREDIT')
            )
        ]
    )
)
```

In the above code, we list out the first five accounts on debit, in descending order of the value, creating the column ZF_BSEG_HKONT_DEBIT. We also list out the first ten accounts on credit, in descending order of the value, creating the column ZF_BSEG_HKONT_CREDIT.

Dataframe: sort/ sort([])

Sort a dataframe – one column: sort()

We can sort an entire dataframe. For example, we might want to sort the material movements on input date, in order to see the total value in stock over time: (below we assume that the date is in text format YYYYMMDD so that it will sort properly):

```
ZV_DF_MSEGMKPF = (  
    ZV_DF_MSEGMKPF  
    .sort(  
        'MKPF_CPUDT'  
    )  
)
```

Sort a dataframe – many columns: sort(by = [])

If we have many fields that we are sorting by, we use the key word `by=[]`. We can also sort descending, if we want to see the postings in reverse order.

In the below code we include the time-stamp. We also include the document number and item, in order for the result to be "deterministic", meaning always the same.

```
ZV_DF_MSEGMKPF = (  
    ZV_DF_MSEGMKPF  
    .sort(  
        by = [  
            'MSEG_BUKRS',  
            'MSEG_MATNR',  
            'MKPF_CPUDT',  
            'MKPF_CPUTM',  
            'MSEG_MBLNR',  
            'MSEG_ZEILE'  
        ],  
        descending = [False, False, True, True, True, True]  
    )  
)
```

Note: the key words **by** = and **descending** = are al *keyword arguments* (also known as **kwargs**). We write them for clarity.

Dataframe: group/ agg: sort

Sort within an aggregation – same field: sort()

If we do not want to keep all records and only want to keep the last date per material, we can use a group by on the material number with a sort to obtain the last date.

```
ZV_DF_MSEGMKPF = (  
    ZV_DF_MSEGMKPF  
    .group_by(  
        [  
            'MSEG_BUKRS',  
            'MSEG_MATNR'  
        ]  
    )  
    .agg(  
        PI_POLARS.col('MKPF_CPUDT')  
        .sort(descending=True)  
        .first()  
        .alias('ZF_MKPF_CPUDT_LATEST')  
    )  
)
```

Sort within an aggregation – different field: sort_by()

If we also want to keep the value for the last movement, per material, we could sort by a date, time for the value. To be deterministic, we could also include the document number and item.

```
ZV_DF_MSEGMKPF = (  
    ZV_DF_MSEGMKPF  
    .group_by(  
        [  
            'MSEG_BUKRS',  
            'MSEG_MATNR'  
        ]  
    )  
    .agg(  
        PI_POLARS.col('ZF_MSEG_DMBTR_SIGNED')  
        .sort_by(  
            [  
                'MKPF_CPUDT',  
                'MKPF_CPUTM',  
                'MSEG_MBLNR',  
                'MSEG_ZEILE'  
            ],  
            descending = [True, True, True, True]  
        )  
        .first()  
        .alias('ZF_MSEG_DMBTR_SIGNED_LATEST')  
    )  
)
```

Dataframe: group/ agg: maintain_order

Maintain_order for a group_by()

If our table was already sorted, we can use the keyword argument (kwarg) **maintain_order** to keep the order in the **group_by()**:

```
ZV_DF_MSEGMKPF = (  
    ZV_DF_MSEGMKPF  
    .group_by(  
        [  
            'MSEG_BUKRS',  
            'MSEG_MATNR'  
        ],  
        maintain_order=True  
    )  
    .agg(  
        PI_POLARS.col('MKPF_CPUDT')  
        .first()  
        .alias('ZF_MKPF_CPUDT_LATEST')  
    )  
)
```

Dataframe: join: inner/semi/ anti

Join: inner – add columns from another table and filter

We often want to join data together. For example, add header information to detail information, only keep rows that match:

```
ZV_DF_EKPOEKKO = (  
    ZV_DF_EKPO  
    .join(  
        ZV_DF_EKKO,  
        left_on = ['EKPO_EBELN'],  
        right_on = ['EKKO_EBELN'],  
        how = 'inner'  
    )  
)
```

Join: semi – keep rows found in another table

We often want to limit one table on another table. For example, only obtain journal entries created by the payment program and don't keep any columns from the payment program:

```
ZV_DF_BSEGBKPF = (  
    ZV_DF_BSEGBKPF  
    .join(  
        ZV_DF_REGUH,  
        left_on = ['BSEG_BUKRS', 'BSEG_BELNR', 'BKPF_BUDAT'],  
        right_on = ['REGUH_ZBUKR', 'REGUH_VBLNR', 'REGUH_ZALDT'],  
        how = 'semi'  
    )  
)
```

Join: anti – remove rows found in another table

We often want to ignore rows if they are in another table. For example, purchase orders without change document (no approval record).

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .join(  
        ZV_DF_CDPOS,  
        left_on = ['EKPO_EBELN'],  
        right_on = ['CDHDR_OBJECTID'],  
        how = 'anti'  
    )  
)
```

Dataframe: join: left/ full

Join: left – add columns from another table, keep all left

We often want to add information from another table, but not remove any rows from the current table. For example, add the material description, if a description exists:

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .join(  
        ZV_DF_MAKT,  
        left_on = ['EKPO_MATNR'],  
        right_on = ['MAKT_MATNR'],  
        how = 'left'  
    )  
)
```

Join: full – keep all rows from both tables

We often want to keep all rows from both tables, so that we can find in the next step, any rows from either table that did not match, and also keep all of our values for both data sets. For example, when comparing the total per account in the general ledger, to the total per account in the trial balance:

```
ZV_DF_BSEG_HKONT_TOT = (  
    ZV_DF_BSEG_HKONT_TOT  
    .join(  
        ZV_DF_GLT0_SAKNR_TOT,  
        left_on = ['BSEG_BUKRS', 'BSEG_GJAHR', 'BSEG_HKONT'],  
        right_on = ['GLT0_BUKRS', 'GLT0_RYEAR', 'GLT0_RACCT'],  
        how = 'full'  
    )  
)
```

Dataframe: join: cross

Join: cross – multiply one data frame by another

We often want to multiply one dataframe by another dataframe. For example, if we want to have one row for every day/ time in a window, when looking for mass sales within a particular period.

```
ZV_DF_OFFSETS = PI_POLARS.DataFrame(  
    {  
        'ZF_NU_OFFSET': list(range(ZV_NU_OFFSET))  
    }  
)  
  
ZF_VBAPVBAK = (  
    ZF_VBAPVBAK  
    .with_columns(  
        PI_POLARS.col('VBAK_ERDAT')  
        .str.strptime(PI_POLARS.Date, '%Y%m%d', strict=False)  
        .alias('ZF_VBAK_ERDAT_DT')  
    )  
    .join(  
        ZV_DF_OFFSETS,  
        how='cross'  
    )  
    .with_columns(  
        (  
            PI_POLARS.col('ZF_VBAK_ERDAT_DT') -  
            PI_POLARS.col('ZF_NU_OFFSET')  
        )  
        .dt.total_days()  
        .alias('ZF_VBAK_ERDAT_DT_OFFSET')  
    )  
)
```

The above code multiplies all lines in the sales order table, for each integer from 0 up to ZV_NU_DAYS_WINDOW.

The code then calculates an offset date, which is the sales order date minus the number of days in the offset.

This new column can then be used in a `group_by()` to see if there were sales orders within the same time period that add-up to a total sales order value or volume above a particular threshold.

Dataframe: generate a date range

One row per day between two dates

We often need one row per day, for example to left-join a daily FX rate to our postings.

`PI_POLARS.date_range()` produces a series of dates that can be used as a dataframe column. `eager=True` returns the values immediately instead of as a lazy expression. The values are returned in the form of a Polars series, which can then be included in the dictionary expression for the function `DataFrame()`.

```
ZV_SE_DATES = PI_POLARS.date_range(  
    start=ZV_DT_BEG_DATE,  
    end=ZV_DT_END_DATE,  
    interval='1d',  
    eager=True  
)  
  
ZV_DF = PI_POLARS.DataFrame(  
    {  
        'ZF_DT_DAY': ZV_SE_DATES  
    }  
)
```

Dataframe: unique

Ensure that the key is unique

Before joining two tables, we often want to make sure that the key in the right-hand table is unique. Otherwise, the rows in the left-hand table will be duplicated (except in the case of a semi join). For example, get only one description per material, before adding the description to the list of material movements.

```
ZV_DF_MSEG = (  
  ZV_DF_MSEG  
  .join(  
    ZV_DF_MAKT = (  
      ZV_DF_MAKT  
      .unique(  
        subset=['MAKT_MATNR']  
      )  
    ),  
    left_on = 'MSEG_MATNR',  
    right_on = 'MAKT_MATNR',  
    how = 'left'  
  )  
)
```

Note: in the above code, you would also want to make sure that the description is filtered on English language.

Dataframe: pipe

Pipe a dataframe through a shared function

We often want to apply a function from `Z_SHARED_FUNCTIONS` to a dataframe, while keeping the dot-chain style. Use the Polars method `.pipe()` to pass the Polars dataframe or expression - to the left of the `.pipe()` - as the first argument of the function, along with any extra arguments.

```
ZV_DF = (  
    ZV_DF  
    .pipe(  
        FC_MERGE_WITH_NULL_HANDLING,  
        ZV_DF_OFAC,  
        'ZF_KY_JN_OFAC_LFA1'  
    )  
)
```

In the above example, our custom function `FC_MERGE_WITH_NULL_HANDLING` takes three variables: `ZV_DF`, `ZV_DF_OFAC` and `'ZF_KY_JN_OFAC_LFA1'`.

Dataframe: iterate rows

Loop through each row as a dictionary

We rarely need to iterate rows, because Polars is much faster with vectorised operations. But for row-level logic that cannot be vectorised (for example, expanding a SAP authorization range into individual transaction codes), we use `.iter_rows(named=True)` to get each row as a dictionary.

```
for ZV_DI_ROW in ZV_DF.iter_rows(named=True):
    ZV_ST_USER = ZV_DI_ROW['USR02_BNAME']
    ZV_ST_FROM = ZV_DI_ROW['UST12_VON']
    ZV_ST_TO   = ZV_DI_ROW['UST12_BIS']
    ...
```

In the above code, we use the `named=True`. `named=True` causes the row to be returned as a dictionary, so that the items in each column can be accessed based on the column name. If we do not use `named = true`, then `iter_rows()` would return a tuple, rather than a dictionary. In this case, we would access the values based on the position in the tuple, which would reflect the position of the column in the original dataframe:

```
for ZV_TU_ROW in ZV_DF.iter_rows():
    ZV_ST_USER = ZV_TU_ROW(0)
    ZV_ST_FROM = ZV_TU_ROW(1)
    ZV_ST_TO   = ZV_TU_ROW(2)
    ...
```

A Tuple is the same as a List, with the exception that a Tuple is not mutable. If an object is not mutable, it means that it cannot be changed after creation.

Dataframe: concatenate

Concatenate a list of dataframes

We often want to create multiple dataframes for different scenarios, and then concatenate the results into one dataframe.

For example, we might want to get the start date that a purchase order was blocked.

However, there are many different block flags for purchase orders. Therefore, we might want to do the same logic of getting the start date for each of these block flags.

Instead of writing out the same code for each block flag, we can loop around a list of block flags and store the block flag name and start date in a dataframe on each iteration of the loop. As the loop iterates, we can store these dataframes in a list of dataframes.

Once our loop finished, we can then append all of our dataframes that we find in our list, into one final dataframe.

```
ZV_LI_DFS = []
for ZV_ST_CDPOS_TABNAME_FNAME in ZV_LI_BLOCK_FLAGS:
    ZV_ST_FLAG_START = f'{ZV_ST_CDPOS_TABNAME_FNAME}_START'
    ZV_DF_CDPOSCDHDR_START = (
        ZV_DF_CDPOSCDHDR
        .filter(
            (
                PI_POLARS.col('ZF_CDPOS_TABNAME_FNAME') ==
                ZV_ST_CDPOS_TABNAME_FNAME
            ) &
            (PI_POLARS.col('CDPOS_VALUE_OLD') == '') &
            (PI_POLARS.col('CDPOS_VALUE_NEW') == 'X')
        )
        .with_columns(
            PI_POLARS.col('CDHDR_UPDATE')
            .alias(f'{ZV_ST_FLAG_START}')
        )
    )
    ZV_LI_DFS.append(ZV_DF_CDPOSCDHDR_START)

ZV_DF_CDPOSCDHDR_START = (
    PI_POLARS.concat(
        ZV_LI_DFS,
        how = 'diagonal_relaxed'
    )
)
```

In the above code, we use the Polars function `concat()` to concatenate the list of dataframes into one dataframe. We use the keyword argument (kwargs) `'diagonal_relaxed'`, which means that extra columns are added if the dataframes do not have the same columns, with NULL values added as placeholders, for the rows coming from dataframes that did not have the extra columns.

Dataframe: unpivot

Convert a pivot table to a vertical table (unpivot)

The totals tables in the ECC version of SAP: trial balance (GLT0), supplier totals (LFC1), customer totals (KNC1), fixed assets totals (ANC1) are in pivot table format: values per month as separate month columns.

We often want to convert them from a pivot table to a vertical table: one month column and one value column.

Having these tables in a vertical format, makes it easier for us to make a line chart showing the balance / values per month:

```
ZV_LI_ST_COL_NAMES = ['GLT0_HSL01', ..., 'GLT0_HSL16']

ZV_DF_MELTED = (
    ZV_DF
    .melt(
        id_vars = [
            'GLT0_BUKRS',
            'GLT0_RYEAR',
            'GLT0_RACCT',
            'GLT0_HSLVT'
        ],
        value_vars = ZV_LI_ST_COL_NAMES,
        variable_name = 'ZF_GLT0_PERIOD',
        value_name = 'ZF_GLT0_HSLXX'
    )
)
```

In the above code we can use the function melt() to convert from a pivot table to a vertical table.

- id_vars: is the list of columns that will not be affected
- value_vars: is the list of columns that are to be converted into two new columns:
 - variable_name: is the name for the new column that will hold the name of the original column
 - value_name: is the name for the new column that will hold the value that was in the original column

Dataframe: add row number

Add a case number / row ID

We often want to give each row a unique number. For example, to label each exception with a case number in the results table that we present in the dashboard.

Use `.with_row_index()`. `offset=1` means that the first row will get the number 1 (instead of 0).

```
ZV_DF = (  
    ZV_DF  
    .with_row_index(  
        name = 'ZF_NU_CASE_NUMBER',  
        offset = 1  
    )  
)
```

Row number inside with_columns()

If we need the row number alongside other new columns inside the same `with_columns()`, we can use `PI_POLARS.arange()` based on the height of the dataframe (height is the number of rows).

```
ZV_DF = (  
    ZV_DF  
    .with_columns(  
        [  
            PI_POLARS.arange(1, ZV_DF.height + 1)  
            .cast(PI_POLARS.Utf8)  
            .alias('ZF_ST_CASE_NUM')  
        ]  
    )  
)
```

Column: implode

List out information

If we want to turn a column into a list, we can use `implode()`. This is similar to `.str.join('_')`, except for the following:

- With `implode()` the items remain separate and additional logic can be applied to them
- `implode()` is not limited to items that are of datatype string

```
ZV_DF_CDPOSCDHDR_IBAN_MULTICHG = (  
    ZV_DF_CDPOSCDHDR_IBAN  
    .sort(  
        [  
            'LFA1_LIFNR', # Supplier number  
            'CDHDR_UDATE' # Change date  
        ]  
    )  
    .group_by('LFA1_LIFNR')  
    .agg(  
        [  
            PI_POLARS.col('CDDHDR_UDATE')  
            .implode()  
            .alias('ZF_LI_CDHDR_UDATE'),  
  
            PI_POLARS.col('CDHDR_OBJECTID') # IBAN  
            .implode()  
            .alias('ZF_LI_CDHDR_OBJECTID'),  
  
            PI_POLARS.col('CDHDR_OBJECTID')  
            .n_unique()  
            .alias('ZF_NU_CDHDR_OBJECTID_COUNTU')  
        ]  
    )  
    .filter(  
        PI_POLARS.col('ZF_NU_CDHDR_OBJECTID_COUNTU') > 1  
    )  
)
```

If we export `ZV_DF_CDPOSCDHDR_IBAN_MULTICHG` to Excel, Excel might replace some of our list columns with null. This is a current known Polars bug. To avoid this we can convert the lists to strings:

`.list.join(' | ')` only works if the items in the list are of type string. If the items are of type date or numeric, then we can use `.list.eval()` to cast them to string: Note: `PI_POLARS.element()` is used to obtain each element of a list.

```
ZV_DF_CDPOSCDHDR_IBAN_MULTICHG = (  
    ZV_DF_CDPOSCDHDR_IBAN_MULTICHG  
    .with_columns(  
        [  
            PI_POLARS.col('ZF_LI_DT_CDHDR_UDATE')  
            .list.eval(PI_POLARS.element().dt.strftime('%Y%m%d'))  
            .list.join(' | ')  
            .alias('ZF_ST_CDHDR_UDATE_LIST')  
        ]  
    )  
)
```

Columns: select/ rename

Select fields

We can limit the number of fields in our dataframe to make it lighter:

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .select(  
        [  
            'EKPO_EBELN',  
            'EKPO_EBELP',  
            'EKPO_MATNR',  
            'EKPO_NETWR',  
            'EKPO_NETPR'  
        ]  
    )  
)
```

Rename fields

Sometimes we need to rename fields:

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .rename(  
        {  
            'EBELN': 'EKPO_EBELN',  
            'EBELP': 'EKPO_EBELP',  
            'MATNR': 'EKPO_MATNR',  
            'NETWR': 'EKPO_NETWR',  
            'NETPR': 'EKPO_NETPR'  
        }  
    )  
)
```

Columns: drop

Drop one or more columns

We can drop columns that we no longer need. This keeps the dataframe light and makes the output easier to review. We can also drop columns to avoid confusion when joining multiple tables-

Use a list of column names with `.drop()`. This is the opposite of `.select()`: instead of saying which columns to keep, we say which columns to remove.

```
ZV_DF = (  
  ZV_DF  
  .drop(  
    [  
      'BSEG_XREF1',  
      'BSEG_XREF2',  
      'BSEG_XREF3'  
    ]  
  )  
)
```

Columns: slice/ conditional

Keep a part of a column

We often want to create a column that contains only part of the original column:

```
ZV_DF_EKPO = (  
    ZV_DF_EKPO  
    .with_columns(  
        PI_POLARS.col('EKKO_ERDAT')  
        .str.slice(0, 4)  
        .alias('ZF_EKKO_ERDAT_YEAR')  
    )  
)
```

Create conditional columns

We often want to create columns, based on criteria:

```
ZV_DF_BSEG = (  
    ZV_DF_BSEG  
    .with_columns(  
        PI_POLARS.when(  
            PI_POLARS.col('BSEG_SHKZG') == 'S'  
        )  
        .then(  
            PI_POLARS.col('BSEG_DMBTR')  
        )  
        .when(  
            PI_POLARS.col('BSEG_SHKZG') == 'H'  
        )  
        .then(  
            PI_POLARS.col('BSEG_DMBTR') * -1  
        )  
        .otherwise(  
            PI_POLARS.lit(0)  
        )  
        .alias('ZF_BSEG_DMBTR_SIGNED')  
    )  
)
```

Column: string contains / starts_with

Filter where column contains a substring

We often want to find rows where a column contains a particular word or pattern. For example, we may want to find any SAP authorization value that contains a wildcard '*'. We can use `.str.contains()`, **which accepts a regex string pattern**.

Note that we use `r'\*'` (a raw string) so that Python does not interpret the backslash. The `\*` tells the regex to look for a literal `*`, not the regex meta-character `*`.

```
ZV_DF = (  
    ZV_DF  
    .filter(  
        PI_POLARS.col('UST12_VON')  
        .str.contains(r'\*')  
    )  
)
```

In the above code, `r'\*'` represents two layers:

- `r'...'` – **Python layer**: `r`: raw string: raw string means **do not** interpret at Python layer:
 - For example, Python will consider `\n` as newline.. and `r\n` as the two characters `\` and `n`.
 - If we do not say `r''`, then Python can accidentally remove our `\`, before we get to the regex layer
- `\*` – **regex layer**: `\`: escape: **do not** interpret at regex layer:
 - For example, regex will consider `\*` as: anything containing `*`
 - If we do not put `\`, then regex will consider `*` as meaning “allow for repeated characters”:
 - For example: `str.contains(r'(ab)\1*')` will match `'ab'`, `'abab'`, `'abababab'`

So to summarize, because `str.contains()` expects a regex string.... we use `r'...'` to mean, **do not interpret at the Python level** and `\` to mean, **do not interpret at the regex level**. This way we can say “does my string contain `*?`”

Since we are just looking for cases that contain `*`, we can also use `find()` or `contains()` with `literal = True`:

```
ZV_DF = (  
    ZV_DF  
    .filter(  
        PI_POLARS.col('UST12_VON')  
        .str.find('*')  
    )  
)
```

```
ZV_DF = (  
    ZV_DF  
    .filter(  
        PI_POLARS.col('UST12_VON')  
        .str.contains('*', literal=True)  
    )  
)
```

Filter where column starts with a character

`starts_with()` checks the beginning of the string. For example, SAP authorization placeholders often start with '\$':

```
ZV_DF = (  
    ZV_DF  
    .filter(  
        PI_POLARS.col('USOBT_LOW')  
        .str.starts_with('$')  
    )  
)
```

Column: regex clean data, fuzzy word matching

Keep only alpha numeric

We often want to tidy up names before we do comparisons. For this we can use regex:

```
ZV_DF = (  
  ZV_DF  
  .with_columns(  
    PI_POLARS.col('SDN_NAME')  
    .str.strip_chars()  
    .str.replace_all(  
      r'^\0-9a-zA-Z',  
      ''  
    )  
    .str.to_uppercase()  
    .alias('ZF_SDN_NAME_CLEAN')  
  )  
)
```

Build a regex expression from a list of keywords

We often want to search for any of many keywords entered by the auditor. For example, suspicious words such as 'cash', 'gift' or 'bonus' in the journal entry text.

We use `PI_RE.escape()` so that special regex characters in the keywords are treated literally (for example, a '.' in a keyword stays a literal '.', not 'any character'). If we have a special character, then `PI_RE.escape` will replace it with `\`. For example, 'bonus?' would become 'bonus\? '. Then, when we pass the string to `str.contains()`, it will search for 'bonus?' and not interpret the '?' as a special regex character.

The escaped key words, shown in the `ZV_LI_KEYWORDS` variable below, are joined with '|', giving 'cash|gift|bonus'. `str.contains()` interprets this as a regex pattern, that means: "Return True if you find the words cash or gift or bonus in the text".

In this way, our filter is matching any of our key words, including any characters that are not alpha-numeric, such as '?'.

```
import re as PI_RE  
  
ZV_LI_KEYWORDS = ['cash', 'gift', 'bonus']  
  
ZV_ST_REGEX = '|'.join(  
  [PI_RE.escape(k) for k in ZV_LI_KEYWORDS]  
)  
  
ZV_DF = (  
  ZV_DF  
  .filter(  
    PI_POLARS.col('BKPFBKTXT')  
    .str.contains(ZV_ST_REGEX)  
  )  
)
```

Columns: string <-> numeric conversion

String from numeric

The most common type of string datatype that we use is Utf8 (which is basically ASCII). For example, if we wanted to do a Benford's law statistical analysis on the leading digits of our numbers in the data set, we might use the following to first convert our numeric field into a string:

```
ZV_DF_BSAK = (  
  ZV_DF_BSAK  
  .with_columns(  
    PI_POLARS.col('BSAK_DMBTR')  
    .cast(  
      PI_POLARS.Utf8,  
      strict=False  
    )  
    .str.slice(0, 2)  
    .alias('ZF_BSAK_DMBTR_2_DIG')  
  )  
)
```

For cast and strptime (and many other functions in Polars), we use strict = False to get a null value, rather than having our program bug, in case the number cannot be interpreted.

Numeric from string

If we have a column that is in string format - for example, just after importing data from a text file - then we can convert it to numeric format, if it is listed in the fields to be converted to numeric format:

```
ZV_DF = (  
  ZV_DF  
  .with_columns(  
    [  
      PI_POLARS.col(ZV_ST_COL)  
      .cast(  
        PI_POLARS.Float64,  
        strict=False  
      )  
      for ZV_ST_COL in ZV_LI_NUMERIC_COLS  
    ]  
  )  
)
```

Columns: string <-> date conversion

String from date

If we have a column in date format, we may wish to convert it to string format, so that we can extract certain information, such as the year or the month, or use the field as a join key.

```
ZV_DF_BSID = (  
  ZV_DF_BSID  
  .with_columns(  
    PI_POLARS.col('ZF_BSID_BUDAT_DT')  
    .dt.strftime('%Y%m%d')  
    .alias('ZF_BSID_BUDAT_ST')  
  )  
)
```

Date from string

As part of our naming convention, we decided to always import dates as strings in the format 'YYYYMMDD'. However, if we want to calculate the difference in days between two dates, we need to convert to date.

```
ZV_DF_BSID = (  
  ZV_DF_BSID  
  .with_columns(  
    PI_POLARS.col('BSID_BUDAT')  
    .str.strptime(  
      PI_POLARS.Date,  
      '%Y%m%d',  
      strict=False  
    )  
    .alias('ZF_BSID_BUDAT_DT')  
  )  
)
```

Always import dates as 'YYYYMMDD'

In our naming convention, we always import date fields as strings in the format 'YYYYMMDD'. This is so that we can be consistent when writing code concerning dates.

If we create date variables, we also always create them as strings in the format 'YYYYMMDD'.

Columns: calculate age

Calculate age

In the below code, we calculate, what Polars calls a Duration: a difference between two dates. We then use the method `.dt.total_days()` to convert this duration into days.

```
from datetime import date as PI_DATE

ZV_DT_TODAY = PI_DATE.today()
ZV_DF_EKKO = (
    ZV_DF_EKKO
    .with_columns(
        PI_POLARS.col('EKKO_BEDAT')
        .str.strptime(
            PI_POLARS.Date,
            '%Y%m%d',
            strict=False
        )
        .alias('ZF_EKKO_BEDAT_DT')
    )
    .with_columns(
        (
            PI_DATE.today() -
            PI_POLARS.col('ZF_EKKO_BEDAT_DT')
        )
        .dt.total_days()
        .alias('ZF_EKKO_BEDAT_DT_AGE')
    )
)
```

In the above code, we convert the Duration into days. However, we could have also converted it into hours, minutes or seconds instead:

```
(
    PI_DATE.today() -
    PI_POLARS.col('ZF_EKKO_BEDAT_DT')
)
.dt.total_hours()
.alias('ZF_EKKO_BEDAT_AGE_HOURS')
```

```
(
    PI_DATE.today() -
    PI_POLARS.col('ZF_EKKO_BEDAT_DT')
)
.dt.total_minutes()
.alias('ZF_EKKO_BEDAT_AGE_MINS')
```

```
(
    PI_DATE.today() -
    PI_POLARS.col('ZF_EKKO_BEDAT_DT')
)
.dt.total_seconds()
.alias('ZF_EKKO_BEDAT_AGE_SECS')
```

Columns: compare rows

Current row minus previous row

We often need to compare to a prior row. We can do this using `.diff()`. For example, if we want to compare the date of our purchase order to the date on the previous row, after sorting.

```
ZV_DF_EKPO_EKKO = (  
    ZV_DF_EKPO_EKKO  
    .sort(  
        [  
            'EKPO_WERKS',  
            'EKKO_LIFNR',  
            'EKPO_MATNR',  
            'ZF_EKKO_BEDAT_DT'  
        ]  
    )  
    .with_columns(  
        PI_POLARS.col('ZF_EKKO_BEDAT_DT')  
        .diff()  
        .over(  
            [  
                'EKPO_WERKS',  
                'EKKO_LIFNR',  
                'EKPO_MATNR'  
            ]  
        )  
        .dt.total_days()  
        .fill_null(0)  
        .alias('ZF_EKKO_BEDAT_DAYS_FROM_PREV')  
    )  
)
```

In the above code: we use `.over([])` to make sure that we do not compare dates between different plants, suppliers, materials.

For the first row for each plant, supplier, material we would get `null()`. We may not want null, because null can disrupt other analysis. Therefore, we can say `.fill_null(0)`, to replace the null with 0.

Columns: compare rows

Reference another row

We often need to reference a previous row. For example, we may wish to know if we have a supplier that appears to be a shell company, because they are sending invoices that follow each other incrementally.

```
ZV_DF_BSAIK = (  
    ZV_DF_BSAIK  
    .with_columns(  
        PI_POLARS.col('BSAIK_XBLNR')  
        .cast(PI_POLARS.Utf8)  
        .str.strip_chars()  
        .str.extract_all(r'(\d+)')  
        .list.last()  
        .cast(PI_POLARS.Int64, strict=False)  
        .alias('ZF_BSAIK_XBLNR_NUM')  
    )  
    .sort(  
        [  
            'BSAIK_BUKRS',  
            'BSAIK_LIFNR',  
            'BSAIK_BLDAT'  
        ]  
    )  
    .with_columns(  
        (  
            PI_POLARS.col('ZF_BSAIK_XBLNR_NUM') -  
            (  
                PI_POLARS.col('ZF_BSAIK_XBLNR_NUM')  
                .shift(1)  
                .over(  
                    [  
                        'BSAIK_BUKRS',  
                        'BSAIK_LIFNR'  
                    ]  
                )  
            )  
        )  
        .alias('ZF_BSAIK_XBLNR_NUM_MINUS_PREV')  
    )  
)
```

In the above code, `str.extract_all()` takes a raw string and interprets it as a regex to provide a list of all the strings that matched the pattern. The regex pattern `\d*` means “Take any strings that are one or more digits”. Then we apply to the list: `.list.last()`, which means: “Only take the last item of the list”.

Here we are assuming that the supplier might have varying invoice patterns, but that it is the last numeric part of the invoice number that is incremented, each time they issue an invoice.

We may wish to try other assumptions, such as “the first numeric part of the invoice number that is incremented”.

Columns: expression

Create an expression from a dictionary

The below code creates an expression based on information in a dictionary:

```
ZV_OB_EXPR_BUCKET = None

for ZV_DI_LIMIT in ZV_LI_DI_APPROVAL_LIMITS:
    ZV_ST_LEVEL = ZV_DI_LIMIT['level']
    ZV_NU_THRESHOLD = ZV_DI_LIMIT['threshold_USD']
    ZV_ST_OPERATOR = ZV_DI_LIMIT['operator']

    if ZV_ST_OPERATOR == '<=':
        ZV_BO_CONDITION = (
            PI_POLARS.col('EKPO_NETWR') <= ZV_NU_THRESHOLD
        )

    elif ZV_ST_OPERATOR == '>':
        ZV_BO_CONDITION = (
            PI_POLARS.col('EKPO_NETWR') > ZV_NU_THRESHOLD
        )

    else:
        continue

    if ZV_OB_EXPR_BUCKET is None:
        ZV_OB_EXPR_BUCKET = (
            PI_POLARS.when(ZV_BO_CONDITION)
            .then(
                PI_POLARS.lit(ZV_ST_LEVEL)
            )
        )

    else:
        ZV_OB_EXPR_BUCKET = (
            ZV_OB_EXPR_BUCKET
            .when(ZV_BO_CONDITION)
            .then(
                PI_POLARS.lit(ZV_ST_LEVEL)
            )
        )

ZV_OB_EXPR_BUCKET = (
    ZV_OB_EXPR_BUCKET
    .otherwise(
        PI_POLARS.lit('OutOfRange')
    )
)
```

Create a column using an expression

The below code creates a column based on the expression above:

```
ZV_DF_EKPOEKKO = (
    ZV_DF_EKPOEKKO
    .with_columns(
        [
            ZV_OB_EXPR_BUCKET
            .alias('ZF_ST_APPROVAL_BUCKET')
        ]
    )
)
```

Columns: concatenate

Concatenate fields together

We often want to concatenate string fields. For example, to provide the description for a code: Assuming we have added the field TSTCT_TTEXT_BKPF to the BSEG/BKPF cube, we can concatenate the transaction code with its description:

```
B01_12_IT_DF_BSEGBKPF_CUBE = (  
  B01_12_IT_DF_BSEGBKPF_CUBE  
  .with_columns(  
    PI_POLARS.concat_str(  
      [  
        PI_POLARS.col('BKPF_TCODE'),  
        PI_POLARS.col('TSTCT_TTEXT_BKPF')  
      ],  
      separator='_  
    )  
    .alias('ZF_BKPFTCODE_TSTCTTTEXT')  
  )  
)
```

Columns: cum_sum(): window

Create a window with cum_sum()

We may wish to know if two purchase orders are within the same window. For example, they are close in days to each other.

If there are purchase orders that are close in date range and that have similar information, then it could be that the purchase orders have been split in order to circumvent delegation of authority controls.

We can use cum_sum() to add up the values from previous rows. Since True is also equal to 1, cum_sum() below increments for each new window.

```
ZV_DF_EKPO_EKKO = (
    ZV_DF_EKPO_EKKO
    .sort(
        [
            'ZF_LIFNRERNAMMATNRWERKSAPPBUCK',
            'ZF_EKKO_BEDAT_DT'
        ]
    )
    .with_columns(
        (
            (
                (
                    PI_POLARS.col('ZF_EKKO_BEDAT_DT').diff()
                    .dt.total_days()
                )
                .fill_null(0)
                .abs()
            ) >= int(ZV_NU_WINDOWDAYS)
        ) |
        (
            PI_POLARS.col('ZF_LIFNRERNAMMATNRWERKSAPPBUCK') !=
            (
                PI_POLARS.col('ZF_LIFNRERNAMMATNRWERKSAPPBUCK')
                .shift(1)
            )
        )
    )
    .cum_sum()
    .alias('ZF_NU_CASE_NUMBER')
)
```

In the above code, .diff() gives us the Duration between the date on the current row compared to the previous row. .dt.total_days() converts that Duration into the number of days. In case it was the first row in the dataframe, .fill_null(0) replaces the null value with 0. If the number of days exceeds the ZV_NU_WINDOWDAYS threshold or the supplier/user/material/plant or approval bucket is diferente, then the result is True... which is also equal to 1. .cum_sum() therefore, adds 1 to the field ZV_NU_CAS_NUMBER.

Meaning the case number is incremented each time there is a difference in days with the previous row that is greater than the threshold or when it is a different supplier (LIFNR)/ user(ERNAM)/ material (MATNR)/ plant (WERKS)/ approval bucket.

Columns: cum_sum(): running total

Reverse running total with cum_sum(), over

We often want to know what the running total is in our dataset. For example, if we want to know the total movements to-date.

In the below example, we use a reverse running total to know the total future stock-out movements. This can be useful to see how much of the stock is used in the future.

For example, we might want to see how much of the stock purchased is used in the future.

```
ZV_DF_MSEG_MKPF_TOT_MATNR_MTH = (  
    ZV_DF_MSEG_MKPF_TOT_MATNR_MTH  
    .sort(  
        [  
            'MSEG_BUKRS',  
            'MSEG_MATNR',  
            'ZF_MKPF_BUDAT_YRMTH'  
        ]  
    )  
    .with_columns(  
        [  
            PI_POLARS.col('ZF_MSEG_DMBTR_CREDIT')  
            .cum_sum(reverse=True)  
            .over(  
                [  
                    'MSEG_BUKRS',  
                    'MSEG_MATNR'  
                ]  
            )  
            .alias('ZF_MSEG_DMBTR_CREDIT_FUTURE')  
        ]  
    )  
)
```

In the above code, to get the reverse running total, we use cum_sum() with reverse = True.

We also use over([]) to only accumulate within the same company (MSEG_BUKRS) and material (MSEG_MATNR).

Column: explode

Create a list of words, then rows from a string field

We often want to get one word from a sentence, phrase, or name. For example, in SAP ECC we have the supplier name in LFA1_NAME1. We often want to compare this name to the names in the list of personnel. We can use explode() to get one row per word in the NAME1 column. We then use explode() to get one row per word in the name from personnel. Then we can use join() with how= 'inner' to filter on matched words.

```
ZV_DF_LFA1 = (  
    ZV_DF_LFA1  
    .with_columns(  
        PI_POLARS.col('LFA1_NAME1')  
        .str.split(' ')  
        .alias('ZF_LI_LFA1_NAME1')  
    )  
    .explode('ZF_LI_LFA1_NAME1')  
    .rename(  
        {'ZF_LI_LFA1_NAME1' : 'ZF_ST_LFA1_NAME1_EXPLODE'}  
    )  
)
```

The above code creates a column that contains all of the words in LFA1_NAME1 in the form of a list. A new column is created using the .alias() method.

The new list column is then exploded: for each item, a new row is created. The column is then renamed, so that it reflects our naming convention.

Column/ variable: zone an string/ number

Column: zone a text column

For columns, we can use the Polars method `.str.zfill(8)` to zone the column.

```
ZV_DF = (  
    ZV_DF  
    .with_columns(  
        PI_POLARS.col('BSEG_HKONT')  
        .str.zfill(8)  
    )  
)
```

Column: zone a numeric column

If our column is a numeric type, we need to first convert it to string type:

```
ZV_DF = (  
    ZV_DF  
    .with_columns(  
        PI_POLARS.col('BSEG_DMBTR')  
        .cast(PI_POLARS.Utf8)  
        .str.zfill(8)  
    )  
)
```

Variable: zone a text item

If our list of accounts are strings, then we can use the built-in Python method for strings: `zfill()`, to format with leading zeros. For example, `zfill(8)`: format with leading zeros and length of 8 characters.

```
ZV_LI_ST_CASH_ACCOUNTS = ['56789', '78923']  
  
ZV_LI_ST_CASH_ACCOUNTS = [  
    ZV_ST_ACCOUNT.zfill(8)  
    for ZV_ST_ACCOUNT in ZV_LI_ST_CASH_ACCOUNTS  
]
```

Variable: zone an integer item

We often want to pad an integer with zeros. For example, if you have a list of general ledger account numbers –for cash, then you might want to use these to filter your general ledger. This is because the general ledger account field in SAP has leading zeros.

If your list of cash accounts, is a list of integers, you would convert it to string and add leading zeros using the built-in string formatting syntax used in f-strings. For example, `f'{...:08}'`: format with leading zeros and length of 8:

```
ZV_LI_NU_CASH_ACCOUNTS = [56789, 78923]  
  
ZV_LI_ST_CASH_ACCOUNTS = [  
    f'{ZV_NU_ACCOUNT:08}'  
    for ZV_NU_ACCOUNT in ZV_LI_NU_CASH_ACCOUNTS  
]
```

Variables - conversions

When converting variables, we can use the built-in Python or plain Python libraries. No Python library installs required.

Variable: string from numeric

We may want to convert a numeric variable to a string.

```
ZV_NU_AMOUNT = 12345.67
ZV_ST_AMOUNT = str(ZV_NU_AMOUNT)
```

Variable: integer from string

We may want to convert a string variable to an integer.

```
ZV_ST_THRESHOLD = '500000'
ZV_NU_THRESHOLD = int(ZV_ST_THRESHOLD)
```

Variable: float from string

We may want to convert a string variable to a float.

```
ZV_ST_THRESHOLD = '500000.67'
ZV_NU_THRESHOLD = float(ZV_ST_THRESHOLD)
```

Variable: string from date

We can use the built-in Python `.strftime()` function to convert a Python date variable to a string. Note that `%Y` means YYYY, as opposed to YY.

```
ZV_ST_END_DATE = ZV_DT_END_DATE.strftime('%Y%m%d')
```

Variable: date from string

We can use the Python datetime library and function `datetime`, which has the function `strftime()` to convert a string variable to a date. Note that `%Y` means YYYY, as opposed to YY.

```
from datetime import datetime as PI_DATETIME

ZV_DT_DATE = (
    PI_DATETIME.strptime(
        ZV_ST_DATE,
        '%Y%m%d'
    )
    .date()
)
```

Variables: Filepaths

Create a file path, that works for Windows or Linux:

Given a variable for folder and file name, create a full path that is correct for both Windows (\) and Linux (/):

When we put \ in a string, Python thinks it is an escape. For example \n: new line, \t: tab, \\ backslash. We can use r" to make our string raw, so that the \ is not interpreted by Python as escape. Note that we must not put the last \ when using raw for file-paths.

```
ZV_ST_RAW_PATH = r'C:\folder'

ZV_ST_FILE_PATH = PI_OS.path.join(
    ZVFCI_ST_FOLDER,
    ZVFCI_ST_SOURCE_FILE
)
```

in the above example, we create the file path with the plain Python os library's path.join() function, in order to avoid differences when switching from Windows to Linux environments.

List – access items/ loop

Access an object in a list

The below code gives us the second item in the list:

```
ZV_OB_MY_OBJECT = ZV_LI_MY_LIST[1]
```

Access all objects in a list one-by-one

We can use a for loop to access objects in a list one-by-one. The below code accesses each dictionary of approval limits in our list of approval limits. The code then accesses each value inside each dictionary:

```
for ZV_DI_LIMIT in ZV_LI_DI_APPROVAL_LIMITS:
    ZV_ST_LEVEL      = ZV_DI_LIMIT['level']
    ZV_NU_THRESHOLD  = ZV_DI_LIMIT['threshold_USD']
    ZV_ST_OPERATOR   = ZV_DI_LIMIT['operator']
    print(
        f'Level: {ZV_ST_LEVEL}, '
        f'Threshold: {ZV_NU_THRESHOLD}, '
        f'Operator: {ZV_ST_OPERATOR}'
    )
```

List – zip - access items in parallel lists

Zip: Simultaneously access information from multiple lists, same index

Sometimes we have more than one list, and we want to access the information from the same position in each list at the same time, so that we can compare the values, or so that we can group logical steps together, rather than repeating code. We can use zip, to put the values from the different lists, side-by-side according to their index (position).

For example, the below code takes two lists of field names: one for purchase order block start date and one for purchase order block end date. The lists are in the same order of block flag type (for example, purchasing block, posting block, etc.).

The zip() puts the lists side-by-side, so that we can access the field names as pairs from within the for loop. We then use the column names as variables to bring up the columns for block start date and block end date... one block flag at a time.

This way, we can compare the block start and block end date to the purchase order date. We use this approach to see if there are any purchase orders between block start and end dates for the same block flag.

```
ZV_LI_DF = []
for ZV_ST_ST, ZV_ST_END in zip(ZV_LI_COL_ST, ZV_LI_COL_END):
    ZV_DF = (
        ZV_DF_EKPO_EKKO
        .filter(
            (
                PI_POLARS.col('EKKO_BEDAT') >=
                PI_POLARS.col(ZV_ST_ST)
            ) &
            (
                PI_POLARS.col('EKKO_BEDAT') <=
                PI_POLARS.col(ZV_ST_END)
            ) &
            (
                PI_POLARS.col(ZV_ST_ST).is_not_null()
            ) &
            (
                PI_POLARS.col(ZV_ST_END).is_not_null()
            )
        )
    )
    ZV_LI_DF.append(ZV_DF)

ZV_DF_EKPO_EKKO_WHEN_BL = (
    PI_POLARS.concat(
        ZV_LI_DF,
        how='vertical'
    )
)
```

The above code creates a dataframe on each iteration of the for loop. Each dataframe is a subset of the original, based on the filter applied. Each dataframe created is stored in a list of dataframes. After the for loop has finished, these dataframes are then all appended together using the concat function. In the concat function we use the keyword argument (kwargs) vertical. Vertical means that Python expects the list of fields for each dataframe created to be the same, which is True here because they are all created from the same source dataframe.

Dictionaries – loop

Access all objects in a dictionary one-by-one

We can use:

for ZV_ST_KEY, ZV_DI in <dictionaryName>.items():
to loop through all dictionaries inside a dictionary.

```
for ZV_ST_KEY, ZV_DI_DEPARTMENT in ZV_DI_COMPANY.items():  
    for ZV_ST_KEY, ZV_DI_EMPLOYEE in ZV_DI_DEPARTMENT.items():  
        ZV_ST_NAME = ZV_DI_EMPLOYEE['name']  
        ZV_ST_ROLE = ZV_DI_EMPLOYEE['role']  
        ZV_NU_SALARY = ZV_DI_EMPLOYEE['salary']
```

The above code accesses each department dictionary one-by-one.

Within each department dictionary, the code then accesses each employee dictionary.

Within each employee dictionary, the code then accesses the values identified by the keys: name, role and salary.

If we are not sure if a key exists, we can use ZV_DI_EMPLOYEE.get('name'). If the key does not exist, this will return None, instead of crashing the program.

Access an object in a dictionary

The below code creates a dictionary that has three levels.

```
ZV_DI_MY_DICTIONARY = {  
    'key': {  
        'key': {  
            'key': myObject  
        }  
    }  
}
```

The below code gives us the object "myObject" that we mention above:

```
ZV_OB_MY_OBJECT = ZV_DI_MY_DICTIONARY['key']['key']['key']
```

Dictionaries to dataframes

Convert dictionaries of lists to a dataframe

We can convert a dictionary to a dataframe. For example, if the dictionary is in the format:

```
ZV_DI_LI_CUSTOMERS = {  
    'Name': ['Claire', 'John'],  
    'Country': ['Portugal', 'UK'],  
}
```

The below code will convert this dictionary to a dataframe with two columns: Name | Country.

```
ZV_DF = PI_POLARS.DataFrame(ZV_DI_LI_CUSTOMERS)
```

Convert lists of nested dictionaries to a dataframe -> structures

If we have a list of nested dictionaries:

```
ZV_LI_DI_CUSTOMERS = [  
    {  
        'Name': 'Claire',  
        'Address': {  
            'City': 'Lisbon',  
            'Country': 'Portugal'  
        }  
    },  
    {  
        'Name': 'John',  
        'Address': {  
            'City': 'London',  
            'Country': 'UK'  
        }  
    }  
]  
ZV_DF = PI_POLARS.DataFrame(ZV_LI_DI_CUSTOMERS)
```

After the above code, we will have a dataframe, with the columns Name | Address. However, the Address column will be a structure. A structure is like a dictionary inside a column item. We can convert the structure into two columns: City | Country:

```
ZV_DF = ZV_DF.unnest('Address')
```

Structures are also sometimes used in dataframes, in order to hold two or more pieces of information together for further processing. See section on Levenshtein where we use structures.

Download URL

Download information from a website

We often want to download information from a website, such as sanctions lists:

```
import requests as PI_REQUESTS
import io as PI_IO

def FC_DOWNLOAD_CSV_FROM_URL(ZVFCI_FILENAME):
    ZV_URL = (
        's://sanctionslistservice.ofac.treas.gov'
        f'/api/download/{ZVFCI_FILENAME}'
    )
    ZV_RESPONSE = PI_REQUESTS.get(ZV_URL)
    ZV_RESPONSE.raise_for_status()
    return PI_IO.StringIO(
        ZV_RESPONSE.content.decode('iso-8859-1')
    )

ZV_DF = PI_Pandas.read_csv(
    FC_DOWNLOAD_CSV_FROM_URL('SDN.csv'),
    delimiter = ',',
    header = None,
    names = [
        'SDN_ENT_NUM',
        'SDN_NAME',
        ...
    ]
)
ZV_DF = PI_POLARS.from_Pandas(ZV_DF)
```

The above code reads the website CSV into a Pandas dataframe. The Pandas dataframe is then converted into a Polars dataframe for further processing. We convert the Pandas dataframe into a Polars dataframe, because Polars can handle large data sets quicker.

Download FX rates

Get daily FX rates from Yahoo Finance

We often need historical FX rates to revalue foreign-currency postings. The yfinance library lets us download daily closing prices for any ticker between two dates.

yfinance returns a Pandas dataframe. We can then loop the rows and build a dictionary that has the date (as a string in YYYYMMDD format) as the key and the closing rate as the value.

progress=False stops yfinance from printing a download progress bar to the terminal.

```
import yfinance as PI_YFINANCE

ZV_DF_DATA = PI_YFINANCE.download(
    'EURUSD=X',
    start = ZV_DT_BEG_DATE,
    end    = ZV_DT_END_DATE,
    interval = '1d',
    progress = False
)

ZV_DI_EXC_RATE = {
    idx.strftime('%Y%m%d'): row['Close']
    for idx, row in ZV_DF_DATA.iterrows()
}
```

In the above code we can see that we use idx.strftime. The result from yFinance is a Pandas dataframe. Pandas dataframe rows all have an index name. In the case of the Pandas dataframe from yFinance, the idx name is the date column. The column Close is the column that contains the rate at the end of the day.

Levenshtein / map elements

Compare text from two columns: levenshtein

We often want to compare data from two columns. For example, we might want to compare two addresses. However, `levenshtein.distance()` can only work on variables, not on columns. Therefore, we put the values on each row in a structure (which is kind of like a mini dictionary in each item of a column) and pass each structure to the function using `.map_elements()`:

```
import Levenshtein as PI_LEVENSHTEIN

def FC_INV_LEVENSHTEIN_DISTANCE(
    ZVFCI_ST_1,
    ZVFCI_ST_2
):
    if ZVFCI_ST_1 is None or ZVFCI_ST_2 is None:
        return None
    return round(
        (1 / (PI_LEVENSHTEIN.distance(
            ZVFCI_ST_1,
            ZVFCI_ST_2
        ) + 1)) * 100,
        2
    )

def FC_MAP_INV_LEVENSHTEIN_DISTANCE(ZVFCI_DI):
    return FC_INV_LEVENSHTEIN_DISTANCE(
        ZVFCI_DI['LFA1_NAME1'],
        ZVFCI_DI['SDN_NAME']
    )

ZV_DF_LFA1_OFACADD = (
    ZV_DF_LFA1ADD
    .with_columns(
        PI_POLARS.struct(
            [
                'LFA1_NAME1',
                'SDN_NAME'
            ]
        )
        .map_elements(
            FC_MAP_INV_LEVENSHTEIN_DISTANCE,
            return_dtype=PI_POLARS.Float64
        )
        .alias('ZF_INV_LEVENSHTEINDIST')
    )
)
```

Isolation Forest

Train an Isolation Forest model on an existing dataset/ score for anomalies

Isolation forest helps us to identify transactions that are unusual compared to the rest of a data set. Isolation Forest looks at the whole row, rather than testing columns individually.

The column ZF_NU_ISF_ANOMALY will have the value 1 for normal rows and -1 for unusual rows. The field ZF_NU_ISF_SCORE gives a score: the more negative, the more unusual the row, compared to the rest of the data set.

In the below example, we train the Isolation Forest on an existing dataset and we score the existing data set at the same time.

If we want to use our trained model on a new dataset, we can put it in a dictionary and pickle it – see next slide. Below, you can see that we use numeric, one-hot and hashed dataframes: this is explained on the next slide.

```
from sklearn.ensemble import IsolationForest as PI_ISOLATIONFOREST

ZV_DF_X = (
    PI_POLARS.concat(
        [
            ZV_DF_NUMERIC,
            ZV_DF_ONEHOT,
            ZV_DF_HASH
        ],
        how='horizontal'
    )
    .fill_null(0)
)
ZV_AR_X = ZV_DF_X.to_numpy()

ZF_OB_ISF_MODEL = PI_ISOLATIONFOREST(
    n_estimators=ZV_NU_ESTIMATORS,
    max_samples=ZV_ST_MAX_SAMPLES,
    contamination=ZV_NU_CONTAMINATION,
    max_features=ZV_NU_MAX_FEATURES,
    random_state=ZV_NU_RANDOM_STATE
)
ZF_OB_ISF_MODEL.fit(ZV_AR_X)
ZV_AR_ANOMALY = ZF_OB_ISF_MODEL.predict(ZV_AR_X)
ZV_AR_SCORE = ZF_OB_ISF_MODEL.score_samples(ZV_AR_X)
ZV_DF_PREDICT = (
    ZV_DF_X
    .with_columns(
        PI_POLARS.Series('ZF_NU_ISF_ANOMALY', ZV_AR_ANOMALY),
        PI_POLARS.Series('ZF_NU_ISF_SCORE', ZV_AR_SCORE)
    )
)
```

In the above code you can see that we convert the polars dataframe to an array, using the numpy method `to_numpy()`. An array is like a list, except that all values are of the same type. Isolation Forest requires an array of numbers.

When creating the Isolation Forest object, we set the parameters: `n_estimators`: number of trees, `max_samples`: how many rows each tree sees; `contamination`: percentage of expected anomalies: affects only the outcome of the `predict()`; `max_features`: percentage of columns used in each tree; `random_state`: random seed – ensures repeatable outcome.

Isolation Forest – one-hot columns

Convert columns to one-hot

In the below code, each column of the Dataframe is converted to up to 51 0/1 dummy columns. the 50 columns are for the 50 most frequent values in the column.

This technique is used for category text columns, such as gender, fiscal year, country. The one-hot columns are then included in the Isolation Forest model, instead of the original text category columns.

```
for ZV_ST_COL in ZV_DF.columns:
    ZV_LI_TOP_COL = (
        ZV_DF
        .group_by(ZV_ST_COL)
        .len()
        .sort('len', descending=True)
        .head(50)
        .select(ZV_ST_COL)
        .to_series()
        .to_list()
    )

    ZV_DF_ONEHOT = (
        ZV_DF
        .with_columns(
            PI_POLARS
            .when(
                PI_POLARS.col(ZV_ST_COL)
                .is_in(ZV_LI_TOP)
            )
            .then(PI_POLARS.col(ZV_ST_COL))
            .otherwise(PI_POLARS.lit('OTHER'))
            .alias(ZV_ST_COL)
        )
        .to_dummies()
    )
```

Isolation Forest –hash columns

Create hashed columns

Isolation Forest requires all columns to be numeric. If we have text columns that are not repeated categories - for example, text descriptions, rather than gender, fiscal year or country - we can hash them using the library Feature hasher.

```
from sklearn.feature_extraction import FeatureHasher as PI_FEATUREHASHER

ZV_LI_LI_TO_HASH= (
    ZV_DF
    .with_columns(
        [
            PI_POLARS.col(ZV_ST_COL).cast(PI_POLARS.Utf8).fill_null('')
            for ZV_ST_COL in ZV_DF.columns
        ]
    )
    .select(
        PI_POLARS.concat_list(ZV_DF.columns)
    )
    .alias('ZF_LI_HASH_INPUT')
    .get_column('ZF_LI_HASH_INPUT')
    .to_list()
)

ZV_OB_HASHER = PI_FEATUREHASHER(
    n_features=32,
    input_type='string'
)

ZV_SP_DF_HASH = ZV_OB_HASHER.transform(ZV_LI_LI_TO_HASH)

ZV_LI_COLUMNS_HASH = [
    f'ZF_HASH_{i:03d}'
    for i in range(ZV_SP_DF_HASH.shape[1])
]

ZV_DF_HASH = PI_POLARS.DataFrame(
    ZV_SP_DF_HASH.toarray(),
    schema= ZV_LI_COLUMNS_HASH
)
```

In the above code, each column is converted to text. The columns are then all concatenated into one column, using `concat_list()`. The resulting column has a list on each row.

Next, when we do `get_column().to_list`, we obtain a list of lists. Our list of lists, can then be passed to the method `transform()`, which gives us back a scipy sparse matrix (`csr_matrix`). This matrix is the hashed version of our columns.

Finally, we use the Polars DataFrame function to convert the array version of the scipy sparse matrix back into a dataframe.

At this point our text description fields have been converted into hashed text fields, which are numeric. We can then pass these hashed fields to Isolation Forest.

Note: FeatureHasher is not reversible. Furthermore, FeatureHasher can result in collision, where more than one value get the same hash value. Therefore, despite being useful for Artificial Intelligence modelling, FeatureHasher is not good for soft encryption of values, such as IBAN numbers.

Isolation Forest – pickle

Pickle – save objects from memory to disk

When using Isolation Forest, we create an Isolation Forest object that is trained on our data-set. Normally, we would train on a “normal” dataset. If we want to then re-use our trained model at a later data, in order to apply it to a new dataset, then we need to store the model.

When we re-use a trained model on a later data-set, we can obtain higher scores if the new data-set is very different to the original data-set.

We can store our model by using the python library pickle.

```
import pickle as PI_PICKLE

ZV_DI_ISF_ARTIFACTS = {
    'Isolation Forest Model': ZF_OB_ISF_MODEL,
    'Hasher object': ZV_OB_HASHER,
}

ZV_ST_DI_ISF_ARTIFACTS_NAME = f'ZV_DI_ISF_ARTIFACTS{ZV_ST_ISF_MODEL_NUMBER}.sav'
PI_PICKLE.dump(ZV_DI_ISF_ARTIFACTS, open(ZV_ST_DI_ISF_ARTIFACTS_NAME, 'wb'))
```

In the above code, 'wb' is used: w: open file for writing, b: write with binar mode

spaCy – classifying names

Classify names with spaCy

We often want to classify names. For example, we may want to classify supplier names as PERSON or ORG. spaCy provides pre-trained, compact models like `en_core_web_sm`. These are designed to be fast, memory-efficient and capable of running on CPUs, making them ideal for high-volume, real-time applications. The library allows for the extraction of entities (PERSON, ORG, GPE, MONEY) from large volumes of text, without requiring heavy GPU resources. spaCy also enables the training and fine-tuning of the models.

The below function creates a spaCy object that uses an existing Natural Language Processing (NLP) model. If the NLP exists at the folder path, this NLP is used. Otherwise, the NLP from spaCy is used. To use the NLP from spaCy, it is necessary to download it to the virtual environment, before running the below code. The spaCy NLP can be installed, by running the following code, once your VSC PowerShell is activated for your venv:

```
python -m spacy download en_core_web_sm
```

In the below code, the field `ZF_ST_SPACY_LABEL` is added to the Dataframe `ZV_DF_SUPPLIERS`. The content of the column will be ORG, if the name of the supplier is classified by the spaCy NLP as being ORG.

```
import spacy as PI_SPACY

ZV_OB_SPACY_NLP_PATH = PI_OS.path.join(
    ZVFCI_ST_SOURCES_FOLDER,
    ZVFCI_ST_NLP_FOLDER
)

def FC_CREATE_SPACY_NLP(ZVFCI_OB_SPACY_NLP_PATH):
    if PI_OS.path.exists(ZVFCI_OB_SPACY_NLP_PATH):
        ZV_OB_SPACY_NLP = PI_SPACY.load(ZVFCI_OB_SPACY_NLP_PATH)
    else:
        ZV_OB_SPACY_NLP = PI_SPACY.load('en_core_web_sm')
    return ZV_OB_SPACY_NLP

ZV_OB_SPACY_NLP = FC_CREATE_SPACY_NLP(ZV_OB_SPACY_NLP_PATH)
ZV_LI_LABELS = []
for ZV_ST_COMPANY_NAME in ZV_DF_SUPPLIERS.get_column('Company_name'):
    ZV_ST_LABEL = None
    ZV_OB_SPACY_DOC = ZV_OB_SPACY_NLP(str(ZV_ST_COMPANY_NAME))
    for ZV_OB_ENT in ZV_OB_SPACY_DOC.ents:
        if ZV_OB_ENT.label_ == 'ORG':
            ZV_ST_LABEL = 'ORG'
            break

    if ZV_ST_LABEL is None:
        if len(ZV_OB_SPACY_DOC.ents) > 0:
            ZV_ST_LABEL = ZV_OB_SPACY_DOC.ents[0].label_
        else:
            ZV_ST_LABEL = 'UNKNOWN'

    ZV_LI_LABELS.append(ZV_ST_LABEL)

ZV_DF_SUPPLIERS = (
    ZV_DF_SUPPLIERS
    .with_columns(
        PI_POLARS.Series(
            'ZF_ST_SPACY_LABEL',
            ZV_LI_LABELS
        )
    )
)
```

spaCy – train the model (1/2)

Train a spaCy model (1/2)

The spaCy model `en_core_web_sm` is a light-weight model that has already been trained. However, despite the model having already been trained, it can still make mistakes. Therefore, we can train it to better guess names in the future, when we see that the result is not as we expected.

When we want to train spaCy, we need to provide training data in the form of tuples: (start position, end position, label). We can use a list of supplier names, that we have manually classified in Excel, to train our model:

```
ZV_OB_TRAINING_PATH = PI_OS.path.join(
    ZVFCI_ST_SOURCES_FOLDER,
    ZVFCI_ST_TRAINING_FILE
)

ZV_DF_TRAINING_EXCEL = PI_POLARS.read_excel(
    source= ZV_OB_TRAINING_PATH
)

ZV_LI_TRAINING_DATA = []

for ZV_DI_ROW in ZV_DF_TRAINING_EXCEL.to_dicts():

    ZV_ST_NAME = str(ZV_DI_ROW['Name'])
    ZV_ST_LABEL = str(ZV_DI_ROW['Label'])

    ZV_LI_TRAINING_DATA.append(
        (
            ZV_ST_NAME,
            {
                'entities': [
                    (
                        0,
                        len(ZV_ST_NAME),
                        ZV_ST_LABEL
                    )
                ]
            }
        )
    )
```

On the next page, we show how to apply the list of training data to train the spaCy model.

spaCy – train the model (2/2)

Train a spaCy model (2/2)

In the below code, we use the ZV_LI_TRAINING_DATA and ZV_OB_SPACY_NER created in the previous slides.

The code takes the list of training data and adds entries in the NER using the spaCy update() method. The ner is then saved to disk. This way, if training was done, the NER will be updated for the next time that the spaCy program is used.

```
from spacy.training import Example as PI_SPACY_EXAMPLE

ZV_NU_ITERATIONS = 20

for ZV_ST_TEXT, ZV_DI_ANNOTATION in ZV_LI_TRAINING_DATA:
    for ZV_NU_START, ZV_NU_END, ZV_ST_LABEL in ZV_DI_ANNOTATION['entities']:
        ZV_OB_SPACY_NLP.get_pipe('ner').add_label(ZV_ST_LABEL)

ZV_LI_OTHER_PIPES = [
    ZV_ST_PIPE
    for ZV_ST_PIPE in ZV_OB_SPACY_NLP.pipe_names
    if ZV_ST_PIPE != 'ner'
]

with ZV_OB_SPACY_NLP.disable_pipes(*ZV_LI_OTHER_PIPES):

    ZV_OB_OPTIMIZER = ZV_OB_SPACY_NLP.resume_training()

    for i in range(ZV_NU_ITERATIONS):
        PI_RANDOM.shuffle(ZV_LI_TRAINING_DATA)
        ZV_DI_LOSSES = {}

        ZV_LI_EXAMPLES = [
            PI_SPACY_EXAMPLE.from_dict(
                ZV_OB_SPACY_NLP.make_doc(ZV_ST_TEXT),
                ZV_DI_ANNOTATION
            )
            for ZV_ST_TEXT, ZV_DI_ANNOTATION in ZV_LI_TRAINING_DATA
        ]

        ZV_OB_SPACY_NLP.update(
            ZV_LI_EXAMPLES,
            sgd=ZV_OB_OPTIMIZER,
            losses=ZV_DI_LOSSES
        )

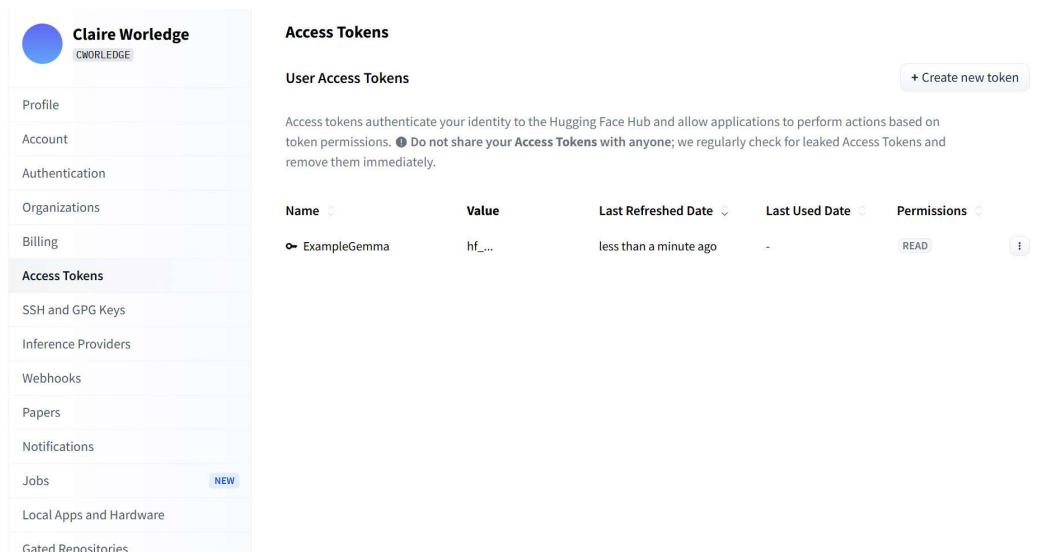
    ZV_OB_SPACY_NLP.to_disk(ZV_OB_SPACY_MODEL_PATH)
```

Hugging Face – login and download a model

Use an API to Hugging Face in order to download a Gemini model

Sometimes we want to use models that are more powerful than the python libraries. On the website <https://huggingface.co>, we can find Artificial Intelligence models that are free to use.

We can create a token on huggingface.co by going to <https://huggingface.co/settings/tokens> and adding a token:



Once we have added a token, we can login to Hugging Face using our token. We can then download a model from HuggingFace. For example, we can download the Gemini model:

```
from huggingface_hub import PI_HUGGING_FACE_LOGIN
import torch as PI_TORCH
from transformers import AutoModelForCausalLM as PI_AUTOMODELFORCAUSALLM

PI_HUGGING_FACE_LOGIN('hf_XXXXJsPbaZ')

ZV_ST_MODEL_ID = 'google/gemma-3-1b-it'

ZV_OB_GEMINI_MODEL = PI_AUTOMODELFORCAUSALLM.from_pretrained(
    ZV_ST_MODEL_ID,
    device_map='auto',
    torch_dtype=PI_TORCH.float16
)
```

in the above code, we can see that from pyTorch we are loading the model with float16. This makes the model lighter, but less accurate, than if we used float32. When we use AutoModelForCausalLM, we are using a text-generation model. A causal model predicts the next tokens (or words) based on previous words. This is why it is used for chatbots, text generation, etc. This library creates a pytorch model, which is also why we use torch.

Gemini model API

Use a Gemini model to generate an output

In the below code we prompt the Gemini model, created in the previous slide. This is like using ChatGPT, but locally and for free. In the below example, we query the Gemini model to ask it if it can find matching meaning for some categories, from within a string content. For example, we could use the below approach to search for key meaning phrases within contracts, in order to check that all required clauses are present.

```
def FC_GENERATE_RESPONSE(ZVFCI_ST_CONTENT, ZVFCI_LI_CATEGORIES):
    ZV_DI_OUTPUTS = {}
    for ZV_ST_CATEGORY in ZVFCI_LI_CATEGORIES:
        ZV_ST_PROMPT = (
            f"Extract only the exact information about '{ZV_ST_CATEGORY}' from the following content. "
            "Do not explain, summarize, or rephrase. If not found, answer 'None'.\n\nContent:\n"
            + ZVFCI_ST_CONTENT
        )

        ZV_LI_DI_CHAT = [
            {'role': 'user', 'content': ZV_ST_PROMPT}
        ]

        ZV_ST_CHAT_TEMPLATE = ZV_OB_TOKENIZER.apply_chat_template(
            ZV_LI_DI_CHAT,
            tokenize=False,
            add_generation_prompt=True
        )

        ZV_LI_LI_INPUT_TOKEN_IDS = ZV_OB_TOKENIZER.encode(
            ZV_ST_CHAT_TEMPLATE,
            add_special_tokens=False,
            return_tensors='pt'
        ).to(ZV_OB_GEMINI_MODEL.device)

        ZV_LI_LI_INPUT_OUTPUT_TOKEN_IDS = ZV_OB_GEMINI_MODEL.generate(
            input_ids=ZV_LI_LI_INPUT_TOKEN_IDS,
            do_sample=True,
            temperature=0.3,
            top_p=0.9,
            max_new_tokens=512
        )

        ZV_ST_OUTPUT = ZV_OB_TOKENIZER.decode(
            ZV_LI_LI_INPUT_OUTPUT_TOKEN_IDS[0][ZV_LI_LI_INPUT_TOKEN_IDS.shape[-1]:],
            skip_special_tokens=True
        ).strip()

        ZV_DI_OUTPUTS[ZV_ST_CATEGORY] = ZV_ST_OUTPUT

    return ZV_DI_OUTPUTS
```

We can see in the above code, that we create a chat string, that is then tokenized and converted to tensor (converted to integers). The tensored version is passed to the Gemini model that passes back a tensor output. The tensor output is then decoded back into a string. The resulting string is added to the dictionary of outputs.

OpenAI API (1/2)

Automatically query open AI based on data analytics results (1/2)

You may be familiar with talking to openAI. We can automate conversations with openAI, in order to obtain feedback on data analytics results. In the prompt, that we create below, we can include data that was presented graphically. We convert the data to a dictionary and then to a JSON. Normally we would also import a JSON to FC_CREATE_AI_PROMPT, with instructions for the particular data-set. In this way, we can allow openai to benefit from our experience/ knowledge/ information about the context relating to the data.

```
from openai import OpenAI as PI_OPENAI
import json as PI_JSON
from openai import AuthenticationError as PI_AUTHENTICATION_ERROR

from dotenv import load_dotenv as PI_LOAD_DOTENV
PI_LOAD_DOTENV(override=True)
ZV_ST_CHATGPT_KEY = PI_OS.getenv('ZV_ST_CHATGPT_KEY')

def FC_CREATE_AI_CLIENT():
    if ZV_AI_CLIENT is None:
        ZV_AI_CLIENT = PI_OPENAI(
            default_headers={"OpenAI-Beta": "assistants=v2"},
            api_key=ZV_ST_CHATGPT_KEY
        )

    ZV_THREAD = ZV_AI_CLIENT.beta.threads.create()
    ZV_THREAD_ID = ZV_THREAD.id
    return ZV_AI_CLIENT, ZV_THREAD_ID

def FC_CREATE_AI_ASSISTANT(ZVFCI_AI_CLIENT):
    ZV_ASSISTANT = ZVFCI_AI_CLIENT.beta.assistants.create(
        name="SAP Dashboard Risk Analyst",
        instructions="""
        You are a skilled Data Analyst analyzing SAP dashboards (Order to Cash and others). Your tasks include:
        1. **Observations**:
        - Review charts for patterns, anomalies, or inefficiencies.
        ...
        - Example:
        """
        | Issue      | Priority | Observation      | Risk              | Recommendation      |
        |-----|-----|-----|-----|-----|
        | Excess Billing | High    | Billing > Orders | Revenue inflation risk | Review periods of variance |
        """
        model="gpt-4o-mini",
        tools=[]
    )
    return ZV_ASSISTANT

def FC_CREATE_AI_PROMPT(
    ZVFCI_DF
):
    ZV_DI_DATA = ZVFCI_DF.to_dict(as_series=False)

    ZV_ST_PROMPT = (
        f"Analysis Task: ....}\n"
        f"Data:\n{PI_JSON.dumps(ZV_DI_DATA, indent=2)}\n\n"
        f"Please provide your response with:\n"
        f"1. ### Observations section\n"
        f"2. ### Risk Assessment section (as a markdown table)\n"
        f"3. ### Executive Summary section\n"
    )
    return ZV_ST_PROMPT
```

OpenAI API (2/2)

Automatically query open AI based on data analytics results (2/2)

In the below code we call openAI with our client, assistant and prompt, that we created on the previous page:

```
def FC_ASK_OPENAI(ZVFCI_AI_CLIENT, ZVFCI_AI_ASSISTANT, ZVFCI_AI_THREAD_ID, ZVFCI_ST_PROMPT):
```

```
    ZVFCI_AI_CLIENT.beta.threads.messages.create(
        thread_id=ZVFCI_AI_THREAD_ID,
        role="user",
        content=ZVFCI_ST_PROMPT
    )
```

```
    try:
```

```
        ZV_AI_RUN = ZVFCI_AI_CLIENT.beta.threads.runs.create_and_poll(
            thread_id=ZVFCI_AI_THREAD_ID,
            assistant_id=ZVFCI_AI_ASSISTANT.id,
            tools=[]
        )
```

```
    except PyOpenAIAuthenticationError as e:
```

```
        print("Authentication failed:", str(e))
```

```
    ZV_ST_RESPONSE_MESSAGE = "No response found."
```

```
    ZV_LI_RESPONSE_MESSAGES=[]
```

```
    if ZV_AI_RUN:
```

```
        ZV_OB_LI_AI_MSG = ZVFCI_AI_CLIENT.beta.threads.messages.list(
            thread_id=ZVFCI_AI_THREAD_ID,
            run_id=ZV_AI_RUN.id
        )
```

```
        for ZV_OB_AI_MSG in ZV_OB_LI_AI_MSG.data:
```

```
            if ZV_OB_AI_MSG.role == "assistant":
```

```
                for ZV_OB_CONTENT in ZV_OB_AI_MSG.content:
```

```
                    if hasattr(ZV_OB_CONTENT, 'text'):
```

```
                        ZV_RESPONSE_MESSAGES.append(
```

```
                            ZV_OB_CONTENT.text.value
```

```
                        )
```

```
    return ZV_LI_RESPONSES
```

```
ZV_AI_CLIENT, ZV_AI_THREAD_ID = FC_CREATE_AI_CLIENT()
```

```
ZV_AI_ASSISTANT = FC_CREATE_AI_ASSISTANT(ZV_AI_CLIENT)
```

```
ZV_ST_PROMPT = FC_CREATE_AI_PROMPT(ZV_DF)
```

```
ZV_LI_RESPONSES = FC_ASK_OPENAI(
```

```
    ZV_AI_CLIENT,
```

```
    ZV_AI_ASSISTANT,
```

```
    ZV_ST_PROMPT
```

```
)
```

The above code creates a list of messages, that we can then output to our end user or use in an audit report.

easyocr: read images (1/2)

Extract data from images (1/2)

Quite often, auditors want to be able to read images, such as Travel & Expense receipts. For example, to check if the name on the receipt matches the employee.

```
import easyocr as PI_EASYOCR

ZV_OB_READER = PI_EASYOCR.Reader(['en'], gpu=False)

ZV_LI_OCR_RESULTS = ZV_OB_READER.readtext(
    ZV_ST_IMAGE_PATH,
    detail=1,
    mag_ratio=2.0
)

def FC_OCR_MERGE_LINES_WITH_SEGMENTS(ZVFCI_LI_OCR_RESULTS, ZVFCI_NU_Y_THRESHOLD, ZVFCI_NU_X_GAP_THRESHOLD):
    ZV_LI_LINES = []

    for ZV_LI_BBOX, ZV_ST_TEXT, ZV_NU_CONF in ZVFCI_LI_OCR_RESULTS:
        ZV_NU_Y_TOP = min(ZV_LI_POINT[1] for ZV_LI_POINT in ZV_LI_BBOX)
        ZV_NU_X_LEFT = min(ZV_LI_POINT[0] for ZV_LI_POINT in ZV_LI_BBOX)
        ZV_NU_X_RIGHT = max(ZV_LI_POINT[0] for ZV_LI_POINT in ZV_LI_BBOX)

        ZV_ST_TEXT = ZV_ST_TEXT.strip()
        for ZV_DI_LINE in ZV_LI_LINES:
            if abs(ZV_DI_LINE['y_top'] - ZV_NU_Y_TOP) < ZVFCI_NU_Y_THRESHOLD:
                ZV_DI_LINE['items'].append({
                    'text': ZV_ST_TEXT,
                    'x_left': ZV_NU_X_LEFT,
                    'x_right': ZV_NU_X_RIGHT
                })
                break
        else:
            ZV_LI_LINES.append({
                'y_top': ZV_NU_Y_TOP,
                'items': [{
                    'text': ZV_ST_TEXT,
                    'x_left': ZV_NU_X_LEFT,
                    'x_right': ZV_NU_X_RIGHT
                }]
            })
    ZV_LI_LINES.sort(key=lambda ZV_DI_L: ZV_DI_L['y_top'])

    ZV_LI_DI_MERGED = []
    for ZV_DI_LINE in ZV_LI_LINES:
        ZV_LI_ITEMS = sorted(ZV_DI_LINE['items'], key=lambda ZV_DI_I: ZV_DI_I['x_left'])
        ZV_LI_SEGMENTS = []
        if not ZV_LI_ITEMS:
            continue
        ZV_DI_CURRENT = ZV_LI_ITEMS[0]
        for ZV_DI_ITEM in ZV_LI_ITEMS[1:]:
            if ZV_DI_ITEM['x_left'] - ZV_DI_CURRENT['x_right'] > ZVFCI_NU_X_GAP_THRESHOLD:
                ZV_LI_SEGMENTS.append(ZV_DI_CURRENT)
                ZV_DI_CURRENT = ZV_DI_ITEM
            else:
                ZV_DI_CURRENT = {
                    'text': ZV_DI_CURRENT['text'] + ' ' + ZV_DI_ITEM['text'],
                    'x_left': ZV_DI_CURRENT['x_left'],
                    'x_right': ZV_DI_ITEM['x_right']
                }
        ZV_LI_SEGMENTS.append(ZV_DI_CURRENT)
        ZV_LI_DI_MERGED.append(
            {
                'y_top': ZV_DI_LINE['y_top'],
                'segments': ZV_LI_SEGMENTS
            }
        )
    return ZV_LI_DI_MERGED
```

easyocr: read images (2/2)

Extract data from images (2/2)

In the below script, we call the function from the previous page and then print out the text from the image. Each segment on the same line is printed, with a space between the segment. Between each line a carriage return is printed.

```
ZV_NU_Y_THRESHOLD = 15
ZV_NU_X_GAP_THRESHOLD = 90

ZV_LI_MERGED_LINES = FC_OCR_MERGE_LINES_WITH_SEGMENTS(
    ZV_LI_OCR_RESULTS,
    ZV_NU_Y_THRESHOLD,
    ZV_NU_X_GAP_THRESHOLD
)

for ZV_DI_LINE in ZV_LI_DI_MERGED_LINES:
    for ZV_DI_SEG in ZV_DI_LINE['segments']:
        print(f" {ZV_DI_SEG['text']} ")
    print(f"\n")
```

Streamlit: Your python project becomes an Intranet site!

We can use a Python library called Streamlit, in order to create a Web-browser User Interface for our Python project. This is particularly useful, if you want to put your Python analytics work inside your Intranet site and allow other team members to interact with and run analysis using your Python code.

Suddenly, you truly can transform the work of your department with automation and Artificial Intelligence.

More information about Streamlit can be found here: <https://streamlit.io/>

In the following pages, we summarize the most useful and basic aspects of Streamlit, to get started.

When using Streamlit – and also when talking to or training AI Agents - it is helpful if you know basic markdown. Markdown defines things like borders, text format, etc. You can find the official documentation of markdown here: <https://daringfireball.net/projects/markdown/>, <https://www.markdownguide.org/> or here for some example games involving markdown: <https://www.markdowntutorial.com/>.

When you create a python project with Streamlit, you can provide access to your project to the rest of your team.

When using Streamlit, if you have access to a corporate GitHub or corporate GitLab, you can put your streamlit application inside the corporate GitHub/ GitLab, using **Stlite** (running the application in the user's browser).

For more robust set-ups and handling of heavy data sets, if you have a server, you can host your Streamlit application on a server. However, it can be a good idea to demonstrate your work in GitLab whilst waiting for a Windows server, or in order to do a quick Proof Of Concept/ Business Use Case for your audit team.

Streamlit: create a Web application

To understand what Streamlit is, the best way is to create the below python script and then run it from the PowerShell terminal in VSC, with the Streamlit command, as mentioned below.

As usual, before running the python script, you need to import the libraries to your virtual environment (and then save the libraries to your requirements.txt file – for best practice):

```
pip install streamlit
pip install numpy
pip install pandas
```

Once you have installed the libraries to the virtual environment, create a python script called FC_MY_STREAMLIT_GRAPH.py, and copy the following content to it:

```
import streamlit as PI_STREAMLIT
import numpy as PI_NUMPY
import pandas as PI_PANDAS

ZV_DF_RANDOM_CHART= PI_PANDAS.DataFrame(
    PI_NUMPY.random.randn(20, 3),
    columns=['a', 'b', 'c']
)

PI_STREAMLIT.line_chart(ZV_DF_RANDOM_CHART)
```

Next, you can launch your Streamlit web application (FC_MY_STREAMLIT_GRAPH.py) by typing the following in the VSC terminal:

```
streamlit run FC_MY_STREAMLIT_GRAPH.py
```

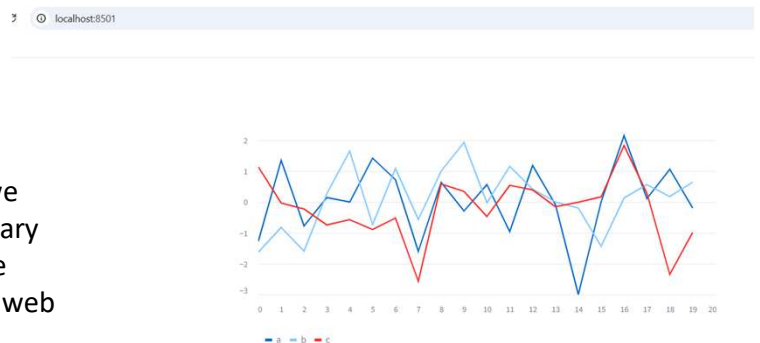
As soon as you hit enter, you will see that a web-browser will appear showing your graph:

You can see from the above code example, that we used the line_chart() function of the Streamlit library to automatically create a graph. We then used the streamlit run command in PowerShell to launch a web application that is running on our local PC.

Streamlit provides you with functions that enable you to automatically generate a wide variety of web-browser user interface widgets, such as interactive graphs, tables, buttons, filters, navigation panes, file upload functions, etc. In this way, you can build an intranet site that takes information from your end-users, that runs your python scripts in the background and that returns the presentation as tables, graphs, filters, etc.

With the Streamlit library installed, you can bring up the Streamlit landing page for more information.

```
streamlit hello
```



Streamlit: design your page with markdown

Render text, labels, dividers or use HTML styling

We can use the `markdown()` function to format text on the page. For example, headings, bold, italic, dividers, and inline HTML.

```
import streamlit as PI_STREAMLIT

# Section headers
PI_STREAMLIT.markdown('### 1. Input & Configuration')
PI_STREAMLIT.markdown('### 2. Processing Results')
PI_STREAMLIT.markdown('### 3. SOX test results and workpaper')

# Bold labels inside a container
PI_STREAMLIT.markdown('**1.1 Upload evidence of system password requirements.**')
PI_STREAMLIT.markdown('**2.2 Overview Results**')

# Divider line
PI_STREAMLIT.markdown('---')

# Italic grey note
PI_STREAMLIT.markdown(
    ".*:grey[**Note:** This table only includes findings for 'Pass' or 'Fail'.]**"
)

# Italic grey note
PI_STREAMLIT.markdown(
    "<b>In-line</b> <span style='color:red;'>HTML</span>",
    unsafe_allow_html=True
)
```

1. Input files

1.1 Upload tab-delimited file.

Note: Note about my table: 'Note'.

In-line HTML

Streamlit: upload files, show table in a box

Container with file upload / container with table

The `FC_FILE_UPLOADER()` function is mentioned in the section on shared functions. This function enables a file to be read into memory by using the streamlit `file_uploader()` function.

In the below example, we call the function `FC_IMPORT_TEXT()`, mentioned in the section on shared functions. This function recognizes that the object is a file in memory (rather than a file path) and returns a dataframe.

The `container()` function draws a box.

We then display the dataframe using the `dataframe()` function with the options:

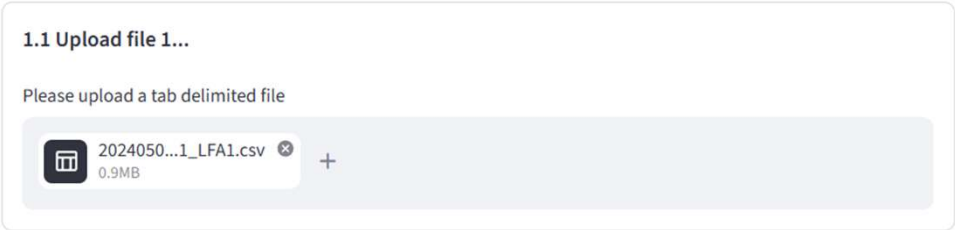
- `hide_index=True`: removes the table index if present
- `width='stretch'` makes the table fill the full container width.

```
from Z_SHARED_FUNCTIONS.FC_IMPORT_TEXT import FC_IMPORT_TEXT

with PI_STREAMLIT.container(border=True):
    PI_STREAMLIT.markdown('**1.1 Upload file 1...**')
    ZV_OB_FILE1 = FC_FILE_UPLOADER(
        ZVFCI_LI_ALLOWED_TYPES: list = ['csv'],
        ZVFCI_BO_ACCEPT_MUL_FILES: = False,
    )

ZV_DF = FC_IMPORT_TEXT(ZV_OB_FILE1)

with PI_STREAMLIT.container(border=True):
    PI_STREAMLIT.markdown('**2.1 Results**')
    PI_STREAMLIT.dataframe(ZV_DF, width='stretch', hide_index=True)
```



After uploading our file, it is shown as a table on the screen.

2.1 Results								
SN	ERDAT	ERNAM	KTOKK	KUNNR	LAND1	LIFNR	LOEVM	NAME1
1	20110405	I009661	0005		CN	IIL-CR-03		CN Truck Carrier
2	20110405	I009661	0005		CN	IIL-CR-04		CN Forwarding Agent
3	20060322	EISENMANN	CPDL			OTA		OTA
4	20111228	I031934	VEND		CN	SA211		Supplier 211 APJ
5	20111228	I031934	VEND		CN	SA212		Supplier 212 APJ

Streamlit: filter tables

Adding a filter to a table

Streamlit does not have built-in table filtering currently. However, we can add table filtering as follows:

```
ZV_ST_FILTER_COLUMN = PI_STREAMLIT.selectbox(
    'Select column to filter',
    ZV_DF.columns
)

ZV_LI_FILTER_VALUES = (
    ZV_DF
    .select(ZV_ST_FILTER_COLUMN)
    .unique()
    .sort(ZV_ST_FILTER_COLUMN)
    .to_series()
    .to_list()
)

ZV_LI_SELECTED_VALUES = PI_STREAMLIT.multiselect(
    'Select values',
    ZV_LI_FILTER_VALUES,
    default=ZV_LI_FILTER_VALUES
)

ZV_DF_FILTERED = (
    ZV_DF
    .filter(
        PI_POLARS.col(ZV_ST_FILTER_COLUMN).is_in(ZV_LI_SELECTED_VALUES)
    )
)

PI_STREAMLIT.dataframe(
    ZV_DF_FILTERED.to_dicts(),
    use_container_width=True
)
```

Streamlit: enter variables

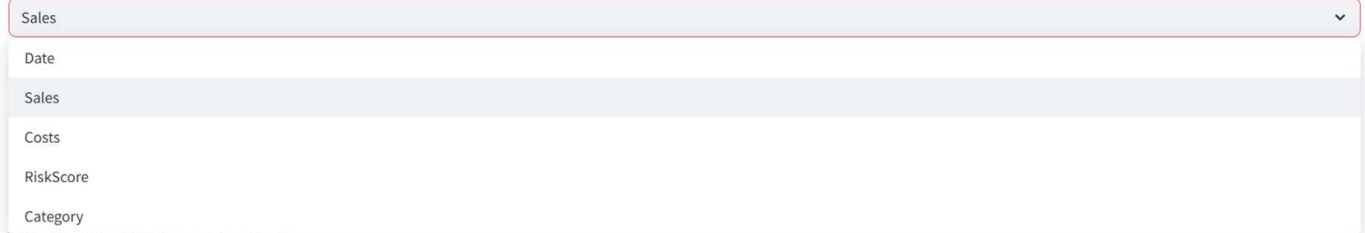
Dropdown list

Once we have uploaded files and created a Dataframe, we can create a dropdown menu based on the values in our Dataframe.

The below code creates a list of columns and then a dropdown box based on the list of columns.

```
ZV_LI_COLUMNS = [  
    ZV_ST_COLUMN  
    for ZV_ST_COLUMN in ZV_DF.columns  
]  
  
if len(ZV_LI_COLUMNS) > 0:  
  
    ZV_ST_SELECTED_NUMERIC_COLUMN = PI_STREAMLIT.selectbox(  
        'Select numeric column',  
        ZV_LI_COLUMNS  
    )
```

Select numeric column



Numeric variable input

For audit tests we often want to include a numeric input: for example, for a threshold value.

```
ZV_NU_THRESHOLD = PI_STREAMLIT.number_input(  
    'Enter threshold value',  
    min_value=0.0,  
    max_value=1000000.0,  
    value=1000.0,  
    step=100.0  
)
```

Multiplier variable

5

Date variable input

For audit tests we often want to include a date input: for example cut-off date.

```
ZV_DT_THRESHOLD = PI_STREAMLIT.date_input(  
    'Enter date threshold',  
    value=PI_DATE.today()  
)
```

Streamlit: enter variables – sliders

Numeric slider

For a more visual effect we can use a slider for variable input:

```
ZV_NU_THRESHOLD_PCT = PI_STREAMLIT.slider(  
    'Threshold (%)',  
    min_value=0,  
    max_value=100,  
    value=10  
)
```

Number of random rows



Date slider

We can use the same streamlit slider function for date ranges:

```
from datetime import date as PI_DATE  
  
ZV_TU_DATE_RANGE = PI_STREAMLIT.slider(  
    'Select date range',  
    min_value=PI_DATE.today(),  
    max_value=PI_DATE.today() + PI_TIMEDELTA(days=365),  
    value=(  
        PI_DATE.today(),  
        PI_DATE.today() + PI_TIMEDELTA(days=30)  
    )  
)  
  
ZV_DF_FILTERED = (  
    ZV_DF  
    .filter(  
        (PI_POLARS.col('Date') >= ZV_TU_DATE_RANGE[0]) &  
        (PI_POLARS.col('Date') <= ZV_TU_DATE_RANGE[1])  
    )  
)
```

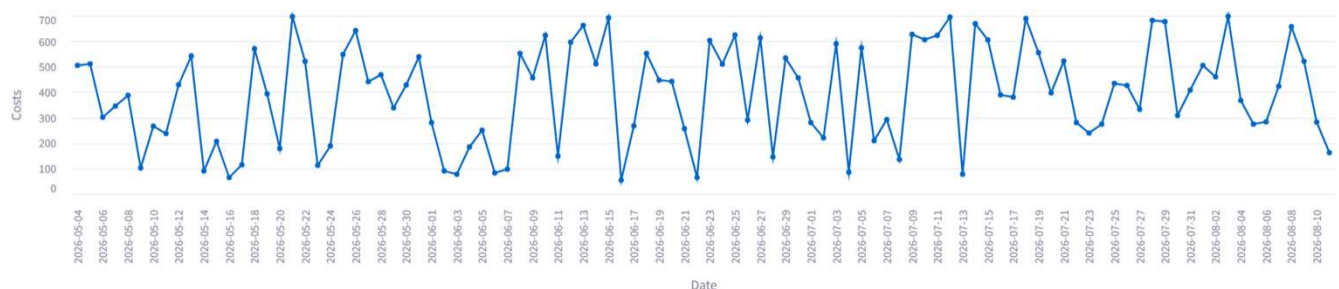
Streamlit: draw graphs

Vega line graph

With Streamlit we can use the Vega Lite graphs to display the data:

```
PI_STREAMLIT.vega_lite_chart(  
    ZV_DF.to_dicts(),  
    {  
        'title': 'Line chart by category',  
        'mark': {'type': 'line', 'point': True, 'tooltip': True},  
        'encoding': {  
            'x': {  
                'field': ZV_ST_SELECTED_CATEGORY_COLUMN,  
                'type': 'nominal'  
            },  
            'y': {  
                'aggregate': 'sum',  
                'field': ZV_ST_SELECTED_NUMERIC_COLUMN,  
                'type': 'quantitative'  
            }  
        }  
    },  
    use_container_width=True  
)
```

Line chart by category



If we add another category field as the colour, then we will get the same graph but split into multiple lines:

```
PI_STREAMLIT.vega_lite_chart(  
    ZV_DF.to_dicts(),  
    {  
        'title': 'Line chart by category',  
        'mark': {'type': 'line', 'point': True, 'tooltip': True},  
        'encoding': {  
            'x': {  
                'field': ZV_ST_SELECTED_CATEGORY_COLUMN,  
                'type': 'nominal'  
            },  
            'y': {  
                'aggregate': 'sum',  
                'field': ZV_ST_SELECTED_NUMERIC_COLUMN,  
                'type': 'quantitative'  
            }  
        },  
        'color': {  
            'field': ZV_ST_SELECTED_CATEGORY_COLUMN2,  
            'type': 'nominal'  
        }  
    },  
    use_container_width=True  
)
```

inspiration for different graphs and documentation can be found here: <https://vega.github.io/vega-lite/examples/>

Streamlit: cross-filtering

Allow filtering on graphs and cross-filtering between graphs

The vega charts can be allowed to filter by using the params keyword. In the below example the type is 'point'. This allows the graph to be clickable. We can allow the graph to be brushable, by using the value 'interval'.

In the below example, the graph is drawn on the page and at the same time stored in ZV_OB_BAR_CHART. the function FC_GET_SELECTION_VALUE() then returns the selected value as variable or a list of values. The function FC_FILTER_BY_CATEGORY_COLUMN() then filters the dataframe. See the section on shared functions for these functions.

In order to allow for cross-filtering on other graphs, we then ensure that all graphs has the option on_select = 'rerun' and that they use data from ZV_DF_FILTERED, rather than ZV_DF.

```
ZV_OB_BAR_CHART = PI_STREAMLIT.vega_lite_chart(
    ZV_DF_FILTERED.to_dicts(),
    {
        'title': 'Bar chart - click category to filter other charts',
        'params': [
            {
                'name': 'ZV_BAR_SELECTION',
                'select': {
                    'type': 'point',
                    'encodings': ['x']
                }
            }
        ],
        'mark': {'type': 'bar', 'tooltip': True},
        'encoding': {
            'x': {'field': ZV_ST_SELECTED_CATEGORY_COLUMN, 'type': 'nominal'},
            'y': {
                'aggregate': 'sum',
                'field': ZV_ST_SELECTED_NUMERIC_COLUMN,
                'type': 'quantitative'
            },
            'color': {
                'condition': {
                    'param': 'ZV_BAR_SELECTION',
                    'field': ZV_ST_SELECTED_CATEGORY_COLUMN,
                    'type': 'nominal'
                },
                'value': 'lightgray'
            }
        }
    },
    use_container_width=True,
    on_select='rerun'
)

ZV_OB_BAR_SELECTION = FC_GET_SELECTION_VALUE(
    ZVFCI_OB_CHART_DATA=ZV_OB_BAR_CHART,
    ZVFCI_ST_PARAM_NAME='ZV_BAR_SELECTION'
)

ZV_DF_FILTERED = FC_FILTER_BY_CATEGORY_SELECTION(
    ZVFCI_DF=ZV_DF,
    ZVFCI_OB_SELECTION=ZV_OB_BAR_SELECTION,
    ZVFCI_ST_CATEGORY_COLUMN=ZV_ST_SELECTED_CATEGORY_COLUMN
)
```

Streamlit: download a file

Download button

The `download_button()` function renders a button that triggers a file download in the browser.

The below script, calls the function `FC_DOWNLOAD_BUTTON`, that we document in in the next section of this document on typical shared functions.

We also document in the typical shared functions section the code for `FC_GET_EXCEL_BYTES`.

```
import streamlit as PI_STREAMLIT
import polars as PI_POLARS

from Z_SHARED_FUNCTIONS.FC_GET_EXCEL_BYTES import FC_GET_EXCEL_BYTES
from Z_SHARED_FUNCTIONS.FC_BTN_DOWNLOAD import FC_BTN_DOWNLOAD

ZV_LI_DI_TABLES = [
    {
        'title': 'Table 1 - Suppliers',
        'filename': 'suppliers',
        'dataframe': PI_POLARS.DataFrame({
            'Supplier': ['A100', 'B200'],
            'Amount': [1000, 2500]
        })
    },
    {
        'title': 'Table 2 - Customers',
        'filename': 'customers',
        'dataframe': PI_POLARS.DataFrame({
            'Customer': ['C100', 'C200'],
            'Amount': [500, 900]
        })
    }
]

for i, ZV_OB_TABLE in enumerate(ZV_LI_DI_TABLES, start=1):

    PI_STREAMLIT.markdown(f"### {ZV_OB_TABLE['title']}")

    ZV_BY_EXCEL_DATA = FC_GET_EXCEL_BYTES(
        ZVFCI_DF=ZV_OB_TABLE['dataframe']
    )

    FC_BTN_DOWNLOAD(
        ZVFCI_BY_DATA=ZV_BY_EXCEL_DATA,
        ZVFCI_ST_FILENAME=ZV_OB_TABLE['filename'],
        ZVFCI_ST_LABEL='Download Excel',
        ZVFCI_ST_KEY=f"download_button_{i}"
    )

    PI_STREAMLIT.dataframe(
        ZV_OB_TABLE['dataframe'],
        width='stretch',
        hide_index=True
    )
```

Streamlit: split the page as columns

Split layout into multiple columns

The function `columns()` divides the screen into side-by-side sections. Pass an integer to split evenly, or a list to control ratios.

the column objects are used in the below code to place two buttons: **Run Analysis** and **Start Over**. The **Run Analysis** and **Start Over** buttons are placed in a 5:1 ratio — **Run Analysis** gets most of the space, **Start Over** stays compact on the right:

```
import streamlit as PI_STREAMLIT

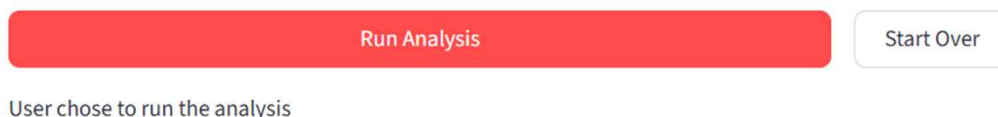
ZV_OB_COL_RUN, ZV_OB_COL_CLEAR = PI_STREAMLIT.columns([5, 1])

with ZV_OB_COL_RUN:
    ZV_BO_RUN_PROCESSING = PI_STREAMLIT.button(
        'Run Processing',
        type='primary',
        width='stretch'
    )

with ZV_OB_COL_CLEAR:
    ZV_BO_CLEAR = PI_STREAMLIT.button(
        'Clear',
        type='secondary',
        width='stretch'
    )

if ZV_BO_RUN_PROCESSING:
    PI_STREAMLIT.write('User chose to run processing')

if ZV_BO_CLEAR:
    PI_STREAMLIT.write('User chose to clear')
```



A 2:0.1:1 ratio can be used to have a middle column with 0.1 as an invisible spacer — it creates a small visual gap between the two sides. the `_` is a Python convention for a variable that is intentionally not used:

```
ZV_OB_COL_LEFT, _, ZV_OB_COL_RIGHT = PI_STREAMLIT.columns([2, 0.1, 1])
```

Streamlit: split the page as tabs

Show content under different tabs

The tabs() function enables you to show content under different tabs.

In the below example, we create two tab objects. Within each tab, a different dataframe is shown.

```
ZV_DF_RANDOM_COST = PI_POLARS.DataFrame(  
    PI_NUMPY.random.randn(20, 3),  
    schema=['EUR', 'USD', 'CNY']  
)  
  
ZV_DF_RANDOM_COST_SUM = (  
    ZV_DF_RANDOM_COST  
    .group_by(  
        PI_POLARS.lit('Total')  
        .alias('')  
    )  
    .agg(  
        PI_POLARS.col('EUR')  
        .sum(),  
  
        PI_POLARS.col('USD')  
        .sum(),  
  
        PI_POLARS.col('CNY')  
        .sum(),  
    )  
)  
  
ZV_OB_TAB_SUMMARY, ZV_OB_TAB_DETAIL = PI_STREAMLIT.tabs(  
    ['Summary', 'Detail']  
)  
  
with ZV_OB_TAB_SUMMARY:  
    PI_STREAMLIT.dataframe(  
        ZV_DF_RANDOM_COST_SUM,  
        width="stretch", hide_index=True  
    )  
  
with ZV_OB_TAB_DETAIL:  
    PI_STREAMLIT.dataframe(  
        ZV_DF_RANDOM_COST,  
        width="stretch", hide_index=True  
    )
```

Summary			
	EUR	USD	CNY
Total	-4.2569	5.6046	6.2061

Detail			
EUR	USD	CNY	
-2.3254	0.9488	0.8809	
-0.3191	0.0195	0.3095	
0.5502	0.4243	-1.1278	
-0.3916	1.0737	0.9289	
2.1969	0.4259	0.0478	
0.5297	0.1753	1.0365	
-0.5596	-0.8925	-0.1555	
-0.9702	2.5028	-1.1555	
0.8497	-0.7601	1.3074	
0.286	0.1951	0.5021	

Streamlit: display an image

Display images

The `image()` function displays images.

The `image()` function accepts a file path, URL, bytes object, or numpy array — making it flexible for both static assets and dynamically processed images.

The `caption` attribute enables a footer for the image.

```
import streamlit as PI_STREAMLIT

# Example image from Aufinia website
ZV_ST_AUFINIA_IMAGE = "https://aufinia.com/wp-content/uploads/2025/07/300FRAMEWORK_Circle.png"

# Display a single image
PI_STREAMLIT.image(
    ZV_ST_AUFINIA_IMAGE,
    caption="Logo 1",
    width=50
)

# Display multiple images in a 2-column grid
ZV_LI_PREVIEW_IMAGES = [
    ('Logo 2', ZV_ST_AUFINIA_IMAGE),
    ('Logo 3', ZV_ST_AUFINIA_IMAGE)
]

ZV_OB_C1, ZV_OB_C2 = PI_STREAMLIT.columns(2)

with ZV_OB_C1:
    ZV_ST_NAME, ZV_OB_IMG = ZV_LI_PREVIEW_IMAGES[0]
    PI_STREAMLIT.image(
        ZV_OB_IMG,
        caption=ZV_ST_NAME,
        width="stretch"
    )

with ZV_OB_C2:
    ZV_ST_NAME, ZV_OB_IMG = ZV_LI_PREVIEW_IMAGES[1]
    PI_STREAMLIT.image(
        ZV_OB_IMG,
        caption=ZV_ST_NAME,
        width="stretch"
    )
```

For example, if you set-up a supplier review application that finds suppliers that are located in residential buildings rather than in office or corporate buildings, then you might want to display the street image in the web page.

Streamlit: success / warning / error / info / progress bar

Give Users Visual Feedback

The functions `success()` and `error()` can be used to confirm uploads and block execution if required inputs are missing. The function `stop()` halts the script:

```
# Confirm uploads
if ZV_LI_UPLOADS:
    PI_STREAMLIT.success(f"Selected {len(ZV_LI_UPLOADS)} file(s)")

if ZV_OB_POLICY_UPLOAD:
    PI_STREAMLIT.success(f"Policy: {ZV_OB_POLICY_UPLOAD.name}")

# Block run if inputs missing
if not ZV_LI_UPLOADS:
    PI_STREAMLIT.error("Missing evidence files.")
    PI_STREAMLIT.stop()
```

The `info()` function shows a message to the end-user. The `progress()` function shows a progress bar.

```
# Warn user about data quality issues
PI_STREAMLIT.warning("3 files could not be parsed and were skipped.")

# Inform user about non-blocking context
PI_STREAMLIT.info("Processing uses the audit period set in the sidebar.")

# Show progress during a long loop
ZV_OB_BAR = PI_STREAMLIT.progress(0)
for i, ZV_OB_FILE in enumerate(ZV_LI_UPLOADS):
    process(ZV_OB_FILE)
    ZV_OB_BAR.progress((i + 1) / len(ZV_LI_UPLOADS))
```

Note: in the above, we use the python built-in function `enumerate()`. `Enumerate()` enables us to loop around a list of objects, obtaining the index position of the object and the object itself at each loop.

Streamlit: session_state()

Remember variables between user interactions

Streamlit reruns the entire script on every user interaction. `session_state` is a dictionary that **persists across reruns** — it lets you remember values, track state, and pass data between different parts of your app.

For example, for each user interaction, we can update the dictionary to remember a particular variable or status.

```
import streamlit as PI_STREAMLIT

ZV_ST_KEY_CONTROL_STATUS = 'control_status'

# Initialize session state
if ZV_ST_KEY_CONTROL_STATUS not in PI_STREAMLIT.session_state:
    PI_STREAMLIT.session_state[ZV_ST_KEY_CONTROL_STATUS] = 'Not_started'

ZV_BO_RUN = PI_STREAMLIT.button('Run Processing')
ZV_BO_CLEAR = PI_STREAMLIT.button('Clear')

if ZV_BO_RUN:
    PI_STREAMLIT.session_state[ZV_ST_KEY_CONTROL_STATUS] = 'Run Processing'

if ZV_BO_CLEAR:
    PI_STREAMLIT.session_state[ZV_ST_KEY_CONTROL_STATUS] = 'not_started'

PI_STREAMLIT.write(
    'Current status:',
    PI_STREAMLIT.session_state[ZV_ST_KEY_CONTROL_STATUS]
)
```

Streamlit: stlite (1/2)

Share your dashboard with others – even if you do not have a server

We can make our Streamlit application available to others by creating an index.html at the root of the project and pushing this to a GitHub or GitLab repository. If you are not familiar with GitHub, see the section on GitHub above.

```
<!doctype html>
<html>
  <head>
    <meta charset="UTF-8" />
    <link rel="stylesheet" href="https://cdn.jsdelivr.net/npm/@stlite/browser@0.85.1/build/stlite.css" />
  </head>

  <body>
    <div id="root"></div>

    <script type="module">
      import { mount } from "https://cdn.jsdelivr.net/npm/@stlite/browser@0.85.1/build/stlite.js";

      const ZV_ST_APP_CODE = await fetch("./02_PROGRAMMES/FC_FILE_UPLOAD_EXAMPLE.py").then(
        response => response.text()
      );

      const ZV_ST_REQUIREMENTS = await fetch("./02_PROGRAMMES/requirements_stlite.txt").then(
        response => response.text()
      );

      const ZV_LI_REQUIREMENTS = ZV_ST_REQUIREMENTS
        .split("\n")
        .map(line => line.trim())
        .filter(line => line && !line.startsWith("#"));

      const ZV_ST_UPLOADER_CODE = await fetch("./02_PROGRAMMES/Z_SHARED_FUNCTIONS/FC_FILE_UPLOADER.py").then(
        response => response.text()
      );

      const ZV_ST_APP_CONFIG_CODE = "ZV_BO_USE_WIDTH = False\n";

      const ZV_ST_IMPORT_CODE = await fetch("./02_PROGRAMMES/Z_SHARED_FUNCTIONS/FC_IMPORT.py").then(
        response => response.text()
      );

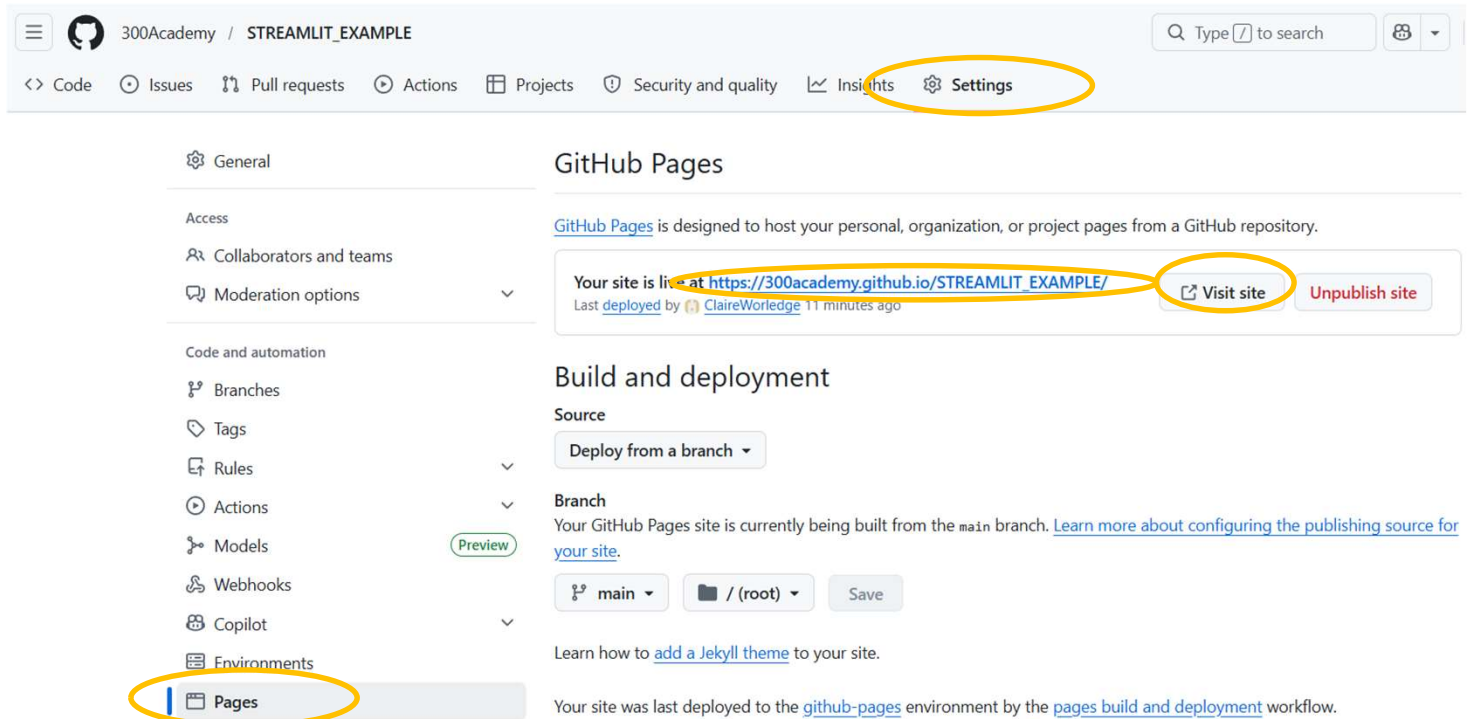
      mount(
        {
          entrypoint: "FC_FILE_UPLOAD_EXAMPLE.py",
          requirements: ZV_LI_REQUIREMENTS,
          files: {
            "FC_FILE_UPLOAD_EXAMPLE.py": ZV_ST_APP_CODE,
            "Z_SHARED_FUNCTIONS/FC_FILE_UPLOADER.py": ZV_ST_UPLOADER_CODE,
            "FC_APP_CONFIG.py": ZV_ST_APP_CONFIG_CODE,
            "Z_SHARED_FUNCTIONS/FC_IMPORT.py": ZV_ST_IMPORT_CODE
          }
        },
        document.getElementById("root")
      );
    </script>
  </body>
</html>
```

The above code uses stlite. The code imports the stlite library, shared functions required for the python to run, the requirements file and any variables. The code then mounts the Streamlit page FC_FILE_UPLOAD_EXAMPLE.py and includes all the files that would normally need to be imported.

Streamlit: stlite (2/2)

View your page from GitHub

Once the project, including your new index.html has been pushed to your corporate GitHub or GitLab, you can then see it by clicking on the link, that you can find in your GitHub Repository -> Settings -> Pages:



If you have put your index.html at the root of your project, you can see the page, by clicking on Visit site.

You will then be able to share the URL to your colleagues in order to share the dashboard centrally.

It is also possible for the StLite to run on results files that are found in folders that are within the same GitHub Repository, by mentioning the folderPath/fileName in the files parameter of the mount() function that we include in index.html.

Typical shared functions

We often do the same task again and again, such as importing a file. We will centralize these tasks in shared functions, so that we only have one copy of the code.

Centralizing tasks in shared functions enables us to only have to maintain it once. Therefore, our task can be continuously improved to the benefit of the entire rest of our project or even our team.

Typical shared functions (basic)

The following code examples give tips for project infrastructure.

These code snippets are typically stored in shared functions that all scripts can access.

They are typically shipped with all projects, to allow for import, export, logging, monitoring, use of external variables, etc.

Start-time, End-time

Record the start and end time of your script

Scripts can take a long time to run, so we may wish to record the start and end times:

```
from datetime import datetime as PI_DATETIME

ZV_ST_STARTTIME = PI_DATETIME.now().strftime(
    '%Y-%m-%d %H:%M:%S'
)

# ...

ZV_ST_ENDTIME = PI_DATETIME.now().strftime(
    '%Y-%m-%d %H:%M:%S'
)

print(
    f'Script started at {ZV_ST_STARTTIME} and ended at {ZV_ST_ENDTIME}'
)
```

Import a text file – columns as String

Import a text file:

When importing text files in Python, we want to make sure that we do not get bugs due to special characters, different encodings or varying delimiters. We can use the below code example to avoid typical import errors. The below script is designed to work with file objects (for example Streamlit updated file objects) or file paths. If the object is a file object, then it will have the `getvalue` attribute. In this case, we can obtain the source by reading the bytes of the file object. The script assumes default delimiter is `'\t'` and tries many default encodings.

```
def FC_IMPORT_TEXT(
    ZVFCI_OB_FILE,
    ZVFCI_ST_DELIMITER=None,
    ZVFCI_LI_ENCODINGS=None
):
    ZV_DF = None
    ZV_ER_LAST = None

    if ZVFCI_ST_DELIMITER == None:
        ZVFCI_ST_DELIMITER = '\t'

    if ZVFCI_LI_ENCODINGS == None:
        ZVFCI_LI_ENCODINGS = [
            'utf8',
            'utf-8-sig',
            'utf8-lossy',
            'windows-1252',
            'latin1',
            'iso-8859-1',
            'iso-8859-15',
            'utf-16',
            'utf-16-le',
            'utf-16-be',
            'cp1250',
            'cp1251',
            'cp1253',
            'cp1254',
            'cp1257'
        ]

    for ZV_ST_ENCODING in ZVFCI_LI_ENCODINGS:
        try:
            if hasattr(ZVFCI_OB_FILE, 'getvalue'):
                ZV_OB_SOURCE = PI_IO.BytesIO(ZVFCI_OB_FILE.getvalue())
            else:
                ZV_OB_SOURCE = ZVFCI_OB_FILE
            ZV_DF = PI_POLARS.read_csv(
                ZV_OB_SOURCE,
                separator=ZVFCI_ST_DELIMITER,
                encoding=ZV_ST_ENCODING,
                ignore_errors=True,
                has_header=True,
                infer_schema_length=0,
                null_values=[''],
                quote_char=None,
                truncate_ragged_lines=True,
                missing_utf8_is_empty_string=True
            ).fill_null('')
            break
        except Exception as ZV_ER:
            ZV_ER_LAST = ZV_ER
    if ZV_DF is None:
        print(f'Last error: {ZV_ER_LAST}')

    return ZV_DF
```

If a delimiter or list of encodings is passed to the above script, then the passed values will be used instead of the default.

Import a text file – columns as Dictionary

Obtain the list of columns in a dataframe

When Python is importing, it can guess the column types. However, these are often, not what we require. Therefore, we can force the column types.

First, we get the list of columns:

```
ZV_LI_COLS = ZV_DF.columns
```

Create a dictionary with the type for each column

Next, we define the type for each column. For example, we can decide that all columns will have the type Utf8.

The below code creates a dictionary that has the key as the field name and the value 'Utf8'.

```
ZV_DI_COLS = {  
    ZV_ST_COL_NAME: PI_POLARS.Utf8  
    for ZV_ST_COL_NAME in ZV_LI_COLS  
}
```

Use a dictionary to force column types

The ZV_DI_COLS can then be used to specify the type for those columns:

```
ZV_DF = PI_POLARS.read_csv(  
    ZV_ST_FILE_PATH,  
    separator = ZVFCI_ST_DELIMITER,  
    encoding = ZV_ST_ENCODING,  
    ignore_errors = True,  
    has_header = True,  
    schema_overrides = ZV_DI_COLS,  
    null_values = [''],  
    quote_char = None,  
    truncate_ragged_lines = True,  
    missing_utf8_is_empty_string = True  
) .fill_null('')
```

note, we can adapt the above approach, by putting the file name and type into a JSON file and then creating a dictionary schema from the JSON file. For example, if we know in advance, that we want certain fields to be of type date or numeric.

See next slide for JSON file import.

Import a json

Import .json to dictionary

The below example uses the built in `open()` Python function to create a file object. The code then uses the `json` library's `load` function to load the file object as a dictionary.

```
import json as PI_JSON

ZV_ST_FILEPATH_APPROVAL_LIMITS = (
    PI_OS.path.join(
        ZF_ST_SOURCES_FOLDER,
        'A_APPROVAL_LIMITS.json'
    )
)

with open(
    ZV_ST_FILEPATH_APPROVAL_LIMITS,
    'r',
    encoding='utf-8'
) as ZV_OB_FILE:
    ZV_DI_APPROVAL_LIMITS = PI_JSON.load(ZV_OB_FILE)
```

In the above code, the built-in **with** keyword is used to enable a file to be opened and then closed once the operations on that file are completed for the open. The keyword **with** can be used in a similar way to allow a connection to a database to be opened... and then close the connection following the completion of the operations that are listed in the tab indent under the `open()`:

.json file structure

A .json file is a text file that:

- has the extension .json, instead of .txt.
- starts and ends with the brackets {}.
- has content written as <keyName>: <object>
- <object> can be of any type:
 - <keyName>: "<text>"
 - <keyName>: [<listItem1, listItem2, listItem3>]
 - <keyName>: <number>
 - <keyName>: {<dictionary>}

In a typical JSON, a keyName like "Approval_limits" might point to an object of type list. Each object in the list is then of type dictionary.

Export to Excel

Export to Excel

We do most of our analysis in Polars. However, to export to Excel, we use a Pandas function.

The Pandas function will try to import the openpyxl library. Therefore, to run the below code, we need to ensure we have openpyxl in our virtual environment. We do not need to import openpyxl to our script because Pandas will import it automatically. As per our naming convention set-up rules, we should openpyxl in our requirements.txt.

```
pip install openpyxl
```

```
def FC_EXPORT_EXCEL(
    ZVFCI_DF_INPUT,
    ZVFCI_ST_RESULTS_FOLDER,
    ZVFCI_ST_RESULTS_FILE
):
    PI_OS.makedirs(
        ZVFCI_ST_RESULTS_FOLDER,
        exist_ok=True
    )
    ZV_ST_FILE_PATH = PI_OS.path.join(
        ZVFCI_ST_RESULTS_FOLDER,
        ZVFCI_ST_RESULTS_FILE
    )
    ZV_DF_Pandas = ZVFCI_DF_INPUT.to_Pandas()
    ZV_DF_Pandas.to_excel(
        ZV_ST_FILE_PATH,
        index = False,
        engine = 'openpyxl'
    )
```

Log memory usage

Log memory usage

Sometimes we want to analyze very heavy data, but our machine is too small and the RAM cannot handle the size. Therefore, we might want to log the RAM usage as the script progresses.

The below script takes the script name and step and outputs to the log the RAM used.

```
import psutil as PI_PSUTIL
import os as PI_OS
from Z_SHARED_FUNCTIONS.FC_LOGGER import ZV_OB_LOGGER

def FC_LOG_MEM_USAGE(ZVFCI_ST_SCRIPT, ZVFCI_ST_SCRIPT_STEP):
    ZV_PROCESS = PI_PSUTIL.Process(PI_OS.getpid())
    ZV_MEM_PROC = ZV_PROCESS.memory_info().rss / (1024 ** 3)
    ZV_MEM_SYS = PI_PSUTIL.virtual_memory()
    ZV_OB_LOGGER.info(
        f'Running: script: {ZVFCI_ST_SCRIPT}, {ZVFCI_ST_SCRIPT_STEP}'
    )
    ZV_OB_LOGGER.info(
        f'[PROCESS] Python process memory: {ZV_MEM_PROC:.2f} GB'
    )
    ZV_OB_LOGGER.info(
        f'[SYSTEM] Total RAM: {ZV_MEM_SYS.total / (1024**3):.2f} GB'
    )
    ZV_OB_LOGGER.info(
        f'[SYSTEM] Used RAM: {ZV_MEM_SYS.used / (1024**3):.2f} GB'
    )
    ZV_OB_LOGGER.info(
        f'[SYSTEM] Available RAM: {ZV_MEM_SYS.available / (1024**3):.2f} GB'
    )
```

Log progress

Log the progress of your script as it runs

If we have heavy data sets our script can take a long time to run. We may wish to log the progress to a log file as the script is running. We can create a logger. In a separate script, called FC_LOGGER, we create:

```
from dotenv import load_dotenv as PI_LOAD_DOTENV
from pathlib import Path as PI_PATH
import os as PI_OS
import logging as PI_LOGGING

PI_LOAD_DOTENV(override=True)

ZV_ST_LOG_FOLDER = PI_OS.getenv('ZV_ST_ETL_FOLDER')
ZV_ST_LOG_MODE = PI_OS.getenv('ZV_ST_LOG_MODE', 'w')
ZV_OB_LOG_FILEPATH = PI_PATH(ZV_ST_LOG_FOLDER) / 'app.log'

ZV_OB_FILEHANDLER = PI_LOGGING.FileHandler(
    filename = ZV_OB_LOG_FILEPATH,
    mode = 'a' if ZV_ST_LOG_MODE == 'a' else 'w',
    encoding = 'utf-8',
    delay = False
)
ZV_OB_FILEHANDLER.setLevel(PI_LOGGING.INFO)
ZV_OB_FILEHANDLER.setFormatter(
    PI_LOGGING.Formatter(
        '%(asctime)s | %(levelname)s | %(message)s'
    )
)

ZV_OB_LOGGER = PI_LOGGING.getLogger(__name__)
ZV_OB_LOGGER.setLevel(PI_LOGGING.INFO)
if not ZV_OB_LOGGER.handlers:
    ZV_OB_LOGGER.addHandler(ZV_OB_FILEHANDLER)
```

Assuming the above script is called FC_LOGGER.py inside our Z_SHARED_FUNCTIONS folder, we can import the logger and then use it to record the progress in the .log file.

```
from Z_SHARED_FUNCTIONS.FC_LOGGER import ZV_OB_LOGGER

ZV_OB_LOGGER.info(
    f'script started at {PI_DATETIME.now().strftime("%Y-%m-%d %H:%M:%S")}'
)
```

Free memory

Release RAM between heavy steps

After a big `group_by()` or `join` we often no longer need certain intermediate dataframes. Even though Python will eventually free the memory, it does not always do this immediately. This can be a problem if our next step is also heavy, because the script may run out of RAM.

We can help Python free the memory by:

- assigning the dataframes to `None` - this removes our reference to them
- calling `PI_GC.collect()` - this asks Python to immediately free the memory of any objects that no longer have references

```
import gc as PI_GC

# release intermediate dataframes
B22_02_TT_DF_USR02      = None
B22_03_TT_DF_UST10C_ALL = None

PI_GC.collect()
```

Despite the above efforts to free-up RAM, we also find that the best way to free RAM is to end the script. This is why it is good practice to have short Python scripts for one task only. Each time a Python script ends, the RAM that it was using will be released.

Alternative environment variable loader

Read Environment Variables Safely

FC_ENV_GET_VALUE resolves a config key by checking two sources in priority order — OS environment first, then a .env file at repo root. Returns an empty string if the key is not found anywhere.

```
def FC_ENV_GET_VALUE(ZVFCI_ST_KEY):
    # Step 1/ Check OS environment first
    ZV_ST_VAL = PI_OS.environ.get(ZVFCI_ST_KEY)
    if ZV_ST_VAL:
        return ZV_ST_VAL

    # Step 2/ Resolve repo root .env
    ZV_OB_ROOT = PI_PATH(__file__).resolve().parents[1]
    ZV_OB_ENV_PATH = ZV_OB_ROOT / ".env"
    if not ZV_OB_ENV_PATH.exists():
        return ""

    # Step 3/ Parse .env file
    with open(ZV_OB_ENV_PATH, "r", encoding="utf-8") as ZV_OB_F:
        for ZV_ST_LINE in ZV_OB_F:
            ZV_ST_LINE = ZV_ST_LINE.strip()
            if not ZV_ST_LINE or ZV_ST_LINE.startswith("#") or "=" not in ZV_ST_LINE:
                continue
            ZV_ST_ENV_KEY, ZV_ST_ENV_VAL = ZV_ST_LINE.split("=", 1)
            if ZV_ST_ENV_KEY.strip() == ZVFCI_ST_KEY:
                return ZV_ST_ENV_VAL.strip()

    return ""
```

Then called like this:

```
ZV_ST_GEMINI_KEY = FC_ENV_GET_VALUE("GEMINI_API_KEY")
ZV_ST_ENV        = FC_ENV_GET_VALUE("APP_ENV")
```

Audit report generation

Generate standardized audit conclusions using structured prompt engineering

We often want to communicate with openAI or Gemini models in order to obtain automated audit report outputs. It helps if we learn markdown to do this. Below is an example of a typical prompt for openAI that uses some very basic markdown.

```
def FC_BUILD_TESTING_RESULTS_PROMPT(
    ZVFCI_ST_CONTROL_DESC: str,
    ZVFCI_ST_TESTING_PROCEDURE: str,
    ZVFCI_ST_EXECUTION_RESULT: str,
) -> str:
    return f"""
    You are a Senior Auditor.

    Based strictly on the information provided below, draft a formal Testing Result &
    Conclusion.

    === CONTROL DESCRIPTION ===
    {ZVFCI_ST_CONTROL_DESC}

    === TESTING PROCEDURES ===
    {ZVFCI_ST_TESTING_PROCEDURE}

    === EXECUTION RESULTS ===
    {ZVFCI_ST_EXECUTION_RESULT}

    Structure your response as follows:

    1. Testing Recap
    Summarize the control objective, audit scope, and testing approach applied.

    2. Results Summary
    Based solely on the Execution Results above:
    - Note any exceptions flagged for review (e.g., access granted before hire date)
    Do NOT fabricate or assume any data not present in the Execution Results.

    3. Conclusion
    State CONTROL PASS or CONTROL FAIL with a one-paragraph justification based on the
    Results Summary and Control Description.

    4. Recommendations
    Provide 2-3 concise, actionable recommendations relevant to the findings. If no exceptions
    were noted, recommend preventive controls and monitoring enhancements.

    Formatting rules:
    - Plain text only, no markdown
    - Use numbers for sections: 1., 2., 3., 4.
    - Use hyphens for bullet points
    - Professional audit tone
    - Maximum 500 words
    """.strip()
```

Streamlit: upload files

Enable the upload of multiple files

The Streamlit `file_uploader()` function gives users a drag-and-drop upload area. `file_uploader()` automatically returns one file object, if one file is selected, or multiple files, if multiple files are selected.

```
import streamlit as PI_STREAMLIT

def FC_FILE_UPLOADER(
    ZVFCI_LI_ALLOWED_TYPES: list = ['csv', 'txt', 'xlsx'],
    ZVFCI_BO_ACCEPT_MUL_FILES: bool = True,
    ZVFCI_ST_KEY: str = 'file_uploader'
):
    return PI_STREAMLIT.file_uploader(
        label='Upload files',
        type=ZVFCI_LI_ALLOWED_TYPES,
        accept_multiple_files=ZVFCI_BO_ACCEPT_MUL_FILES,
        key=ZVFCI_ST_KEY,
        label_visibility='collapsed',
        width='stretch'
    )
```

Note: if the user interacts with the page, by pressing a button, the `file_uploader()` automatically remembers the state for the `file_uploader` object, based on the key name.

To make the streamlit page start over, for example, so that a different file can be uploaded, we can change the name of the key.

This is the reason that we set the key parameter dynamically in the function.

Streamlit: Excel sheet buffer for download

Enable the download of data to Excel

In Streamlit, if we want to download a dataframe to Excel, we first need to store the dataframe as an Excel buffer object. An Excel buffer object stores the data as bytes.

These bytes can then be used to obtain the Excel file.

```
from io import BytesIO
import polars as PI_POLARS

def FC_GET_EXCEL_BYTES(ZVFCI_DF):

    ZV_OB_EXCEL_BUF = BytesIO()

    ZVFCI_DF.write_excel(
        workbook=ZV_OB_EXCEL_BUF,
        worksheet='Results'
    )

    return ZV_OB_EXCEL_BUF.getvalue()
```

In the above code, the dataframe is written into an Excel file that is stored in memory as bytes. The `getvalue()` function returns the value for this Excel file, so that it can then be downloaded.

Streamlit: Excel workbook buffer for download

Enable the download of data to Excel workbook (multiple sheets)

FC_BUILD_EXCEL_BUFFER writes one or more DataFrames into an in-memory Excel file using PI_PANDAS.ExcelWriter. Returns a BytesIO buffer ready to pass directly to a download button — no temp files on disk.

```
def FC_BUILD_EXCEL_BUFFER(ZVFCI_DI_SHEETS: dict) -> PI_BYTES_IO:
    """
    Build Excel file in memory from multiple DataFrames

    Args:
        ZVFCI_DI_SHEETS: dict
            {
                "Sheet Name": DataFrame,
                ...
            }

    Returns:
        BytesIO buffer (ready for download)
    """

    ZV_OB_BUFFER = PI_BYTES_IO()

    with PI_PANDAS.ExcelWriter(ZV_OB_BUFFER, engine="openpyxl") as ZV_OB_WRITER:

        for ZV_ST_SHEET_NAME, ZV_DF_INPUT in ZVFCI_DI_SHEETS.items():

            # Skip invalid DataFrame
            if ZV_DF_INPUT is None:
                continue

            if isinstance(ZV_DF_INPUT, PI_PANDAS.DataFrame) and not ZV_DF_INPUT.empty:
                ZV_DF_INPUT.to_excel(
                    ZV_OB_WRITER,
                    sheet_name=ZV_ST_SHEET_NAME,
                    index=False
                )

    ZV_OB_BUFFER.seek(0)

    return ZV_OB_BUFFER
```

Then called like this:

```
# Single sheet export
ZV_OB_EXCEL_BUF = FC_BUILD_EXCEL_BUFFER({
    "Results": ZV_DF_RESULTS
})

# Multi-sheet export – each DataFrame becomes its own tab
ZV_OB_EXCEL_BUF = FC_BUILD_EXCEL_BUFFER({
    "Summary": ZV_DF_SUMMARY,
    "Details": ZV_DF_DETAILS,
    "Errors": ZV_DF_ERRORS,
})
```

Streamlit: Obtain data from chart selection

Obtain data from chart selection

The function `FC_GET_SELECTION_VALUE()` returns a variable or a list that contains the values selected from a vega chart.

```
def FC_GET_SELECTION_VALUE(ZVFCI_OB_CHART_DATA, ZVFCI_ST_PARAM_NAME):  
    try:  
        return ZVFCI_OB_CHART_DATA.selection[ZVFCI_ST_PARAM_NAME]  
    except Exception:  
        return None
```

Filter Dataframe on selection

The function `FC_FILTER_BY_CATEGORY_SELECTION()`, filters the Dataframe that is used by the graph, other graphs based on the variable or list that is returned by `FC_GET_SELECTION_VALUE()` above. This function therefore helps to support cross-filtering of graphs on a Streamlit page.

```
def FC_FILTER_BY_CATEGORY_SELECTION(  
    ZVFCI_DF,  
    ZVFCI_OB_SELECTION,  
    ZVFCI_ST_CATEGORY_COLUMN  
):  
    if ZVFCI_OB_SELECTION is None:  
        return ZVFCI_DF  
  
    if ZVFCI_ST_CATEGORY_COLUMN not in ZVFCI_OB_SELECTION:  
        return ZVFCI_DF  
  
    ZV_OB_SELECTED_VALUE = ZVFCI_OB_SELECTION[ZVFCI_ST_CATEGORY_COLUMN]  
  
    if isinstance(ZV_OB_SELECTED_VALUE, list):  
        return (  
            ZVFCI_DF  
            .filter(  
                PI_POLARS.col(ZVFCI_ST_CATEGORY_COLUMN).is_in(ZV_OB_SELECTED_VALUE)  
            )  
        )  
  
    return (  
        ZVFCI_DF  
        .filter(  
            PI_POLARS.col(ZVFCI_ST_CATEGORY_COLUMN) == ZV_OB_SELECTED_VALUE  
        )  
    )
```

Streamlit: download button

Button to download files

The `FC_BTN_DOWNLOAD()` function takes data in the form of bytes (for example, as created by the above mentioned `FC_GET_EXCEL_BYTES()` function).

```
import streamlit as PI_STREAMLIT
from datetime import datetime as PI_DATETIME

def FC_BTN_DOWNLOAD(
    ZVFCI_BY_DATA: bytes,
    ZVFCI_ST_FILENAME: str,
    ZVFCI_ST_LABEL: str = 'Download',
    ZVFCI_ST_KEY: str = 'download_button'
):
    ZV_ST_TIMESTAMP = PI_DATETIME.now().strftime('%Y%m%d_%H%M%S')

    return PI_STREAMLIT.download_button(
        label=ZVFCI_ST_LABEL,
        data=ZVFCI_BY_DATA,
        file_name=f'{ZVFCI_ST_FILENAME}_{ZV_ST_TIMESTAMP}.xlsx',
        mime='application/vnd.openxmlformats-officedocument.spreadsheetml.sheet',
        key=ZVFCI_ST_KEY,
        width='stretch'
    )
```

Streamlit: KPIs

Display Metric Cards (useful for Streamlit applications)

This component helps ensure a consistent KPI display across applications.

```
def FC_KPI(
    ZVFCI_ST_TITLE: str,
    ZVFCI_NU_VALUE: int,
    ZVFCI_ST_ICON: str = "",
    ZVFCI_ST_COLOR: str = ""
):
    # Step 1/ Normalize display values
    ZV_ST_VALUE = str(ZVFCI_NU_VALUE)

    # Step 2/ Render KPI card
    PI_STREAMLIT.markdown(
        f"""
        <div class="kpi-card">
            <div class="kpi-icon">{ZVFCI_ST_ICON}</div>
            <div class="kpi-title">{ZVFCI_ST_TITLE}</div>
            <div class="kpi-value {ZVFCI_ST_COLOR}">{ZV_ST_VALUE}</div>
        </div>
        """,
        unsafe_allow_html=True,
    )
```

Then called inside a column layout, one KPI per column:

```
ZV_OB_C1, ZV_OB_C2, ZV_OB_C3 = PI_STREAMLIT.columns(3)

with ZV_OB_C1:
    FC_KPI("# of Screenshots", ZV_NU_SCREENSHOTS, "📷", "blue")

with ZV_OB_C2:
    FC_KPI("# of Records", ZV_NU_RECORDS, "📄", "green")

with ZV_OB_C3:
    FC_KPI("# of Unique Tickets", ZV_NU_TICKETS, "🎫", "amber")
```

Streamlit: session prefix generator

Encapsulates dynamic prefix generation for consistent session state management

Standardize naming conventions across the application, avoids key collisions in session state, enables dynamic scoping based on current context (e.g. test/system), and centralizes prefix logic, making workflows easier to manage and maintain.

```
def FC_GET_PREFIX():  
    """  
    Utility function to get the common prefix for all components and workflows  
    This can be used for consistent naming, logging, or other purposes across the app  
    """  
    return (  
        f"{PI_STREAMLIT.session_state.get('ZV_ST_CURRENT_TEST')}"  
        _f"{PI_STREAMLIT.session_state.get('ZV_ST_CURRENT_SYSTEM')}"  
    )
```

Then called wherever a scoped session key is needed:

```
ZV_ST_PREFIX = FC_GET_PREFIX()  
if f"{ZV_ST_PREFIX}_PROCESSED" not in PI_STREAMLIT.session_state:  
    PI_STREAMLIT.session_state[f"{ZV_ST_PREFIX}_PROCESSED"] = False
```

Dynamically scopes session state keys to prevent conflicts when handling multiple workflows or test scenarios in parallel.

Useful libraries

Python libraries are endless.

They enable us to do statistical Artificial Intelligence to look for rare transactions, read images, read PDFs or Word documents, connect to databases, applications, send results by email or PBI, generate images, text or sound, read data from websites, input information into intranet sites, interpret dashboards, score results for risk, automatically generate audit reports or audit planning...

Here we present some examples of useful libraries. However, in reality there are 100s of thousands of python libraries that are freely available for use.

Python libraries can be found here: [pipy.org](https://pypi.org)

Useful libraries (1/2)

Polars

Handle datasets very fast:

- `to_Pandas()` : convert from Polars to Pandas dataframe
- `read_csv()` : read a csv file

Pandas

Handle datasets:

- `to_excel()` : export to Excel

datetime

Get today's date:

- from datetime import date as PI_DATE
- `PI_DATE.today()`

json

Import a json file to a dictionary:

- `PI_JSON.load()` .. Creates a dictionary out of a file object

rapidfuzz

Fast drop-in replacement for the levenshtein library:

from rapidfuzz.distance import Levenshtein as PI_LEVENSHTEIN. rapidfuzz can also be used to do more advanced scoring for fuzzy matching operations.

levenshtein

Compute the distance between two strings in terms of number of different characters.

requests

Download information from a website.

re

Build and apply regular expressions.

- `escape()` safely escapes user input so that special characters are treated as literals.
- `match()` applies a regex pattern to a string.

yfinance

Download historical FX rates and stock prices from Yahoo Finance. Useful for revaluing foreign-currency postings.

Useful libraries (2/2)

os

Interact with the Operating System:

- `path.join()` : create operating system independent file path
- `makedirs()` : make directory, in case it does not exist

dotenv

- `load_dotenv`: Import secret variables from a `.env` file to RAM
- `dotenv_values`: Import all variables from an end-user variables file

pathlib

- Create paths that are not sensitive to differences between linux and Windows.

logging

- Log the progress of your Python script

gc

Garbage collection.

- `collect()` asks Python to free memory immediately after dropping references to large dataframes.

psutil

Check the RAM used by the Python script