

Artificial Intelligence for Corporate Finance

AI-Augmented Analyst Workflows with Excel, Python and LLMs

Amedeo Andriollo

Finance Department (DRM)
Université Paris Dauphine – PSL

Structure of the course

- **Email:** amedeo.andriollo@dauphine.psl.eu
- **Assessment:** one 150-minute final in-class examination: Part A = 50 minutes; Part B = 100 minutes.
- **Office hours:** feel free to come.
- Slides and materials are updated regularly. Report a typo or mistake: corrections help the whole class.
- **This session:** a first-pass sell-side target study before valuation, with AI as augmentation under evidence and verification discipline.
- **Principal anchors:** Rosenbaum & Pearl, Chapter 1 “Study the Target” / Exhibit 1.3; official OpenAI and Anthropic guidance; Anthropic AI Fluency.

Source basis: Instructor course information; source locators in SLIDE_SOURCE_MAP.md.

The Session 1 professional question

Professional problem

Before a sell-side team values a company, it turns a compact, incomplete pack into a concise view of the business, operating history, unresolved tensions and next questions.

Learning sequence

1. Study TargetCo without a prompting method.
2. Compare outputs and see partial professional feedback.
3. Learn how an LLM and a disciplined workflow change the work.
4. Redo the task in a fresh chat, then inspect the full answer and workflow.

Course path: target-study note → audited Excel workbook/change log → valuation analysis → reconciled Python/Excel analysis → checked deal-extraction table + derived transaction indicators → strategic buy-side briefing/workflow.

Study the target before choosing how to value it

The finance operation

“Study the Target” means learning the company’s story before selecting or benchmarking comparables. A compact target study is neither due diligence nor valuation. It creates the working frame for later valuation choices.

Business profile

- sector and subsector;
- products and services;
- customers and end markets;
- distribution channels;
- geography.

Financial profile – selected Session 1 dimensions

- size;
- profitability;
- growth;
- returns and credit profile later, when data permits.

Source basis: SOURCE-DERIVED: Rosenbaum & Pearl (2013), Ch. 1, Step I “Study the Target” and Exhibit 1.3 Business and Financial Profile Framework.

Reference

Simple measures answer different economic questions

Revenue growth and its disclosed bridge

$$g_t = \frac{R_t}{R_{t-1}} - 1, \quad g_t = v_t + p_t.$$

Here v_t and p_t are disclosed volume and price/mix **contribution points** to revenue growth. **Question answered:** how did sales change, and which disclosed contribution drove the change?

Reported EBITDA margin

$$m_t = \frac{\text{EBITDA}_t}{R_t}.$$

Question answered: how much operating profit is reported per euro of revenue?

Growth and margin can move differently because margin also depends on cost recovery, utilization, mix and other operating factors. **TargetCo:** 2025A utilization exceeds 2022A, while the 2025A reported margin is lower. That is an analytical tension, not proof of causality.

Source basis: SOURCE-DERIVED: Rosenbaum & Pearl (2013), Ch. 1, Exhibit 1.3; TARGETCO APPLICATION.

Reference

Baseline design: the assignment and source pack stay constant

	ChatGPT group	Claude group	Human-only group
Materials	Same two-page TargetCo profile; same VP assignment; same time		
Baseline	Student-composed one-message prompt from common sketch	Student-composed one-message prompt from common sketch	Same finance assignment; no AI
Redo	New Temporary Chat; up to three messages	New Incognito chat; up to three messages	AI access; may split by application
Main comparison	Baseline output/work process versus redo output/work process		

Interpretation discipline

The baseline is not a controlled ChatGPT-versus-Claude model comparison: students may formulate different prompts. It records different student/AI workflows under a common finance assignment, material, time and prompt budget.

Baseline prompt sketch

AI groups: one message, with student freedom

You may edit this modest starting sketch before sending your one permitted baseline message. The VP request remains the separate finance assignment.

Using the supplied TargetCo profile, help me prepare the one-page note requested by the VP. Use the supplied material as the source. Do not value the company.

Human-only group: complete the same finance assignment from the same profile, without AI.

Record: retain the exact prompt your group actually used. Do not add a second message to improve it.

Source basis: INSTRUCTOR SYNTHESIS: student/TargetCo_baseline_prompt_sketch.md.

20-minute baseline: produce a first-pass VP note

Deliverable

One page. Use only the supplied material. Do not value TargetCo. AI groups retain the actual conversation; human-only groups retain notes.

Read for

- business configuration;
- commercial facts and operating history;
- what is known and what is missing.

Do not do

- external research;
- peer selection or valuation;
- a generic diligence inventory;
- a prompting-method exercise.

Work now. We compare outputs at 00:40.

A stronger single-shot instruction

Instructor example: still one message

Use only the supplied TargetCo profile as the permitted source. Prepare one concise, one-page first-pass target-study note for the VP before valuation. Do not use external facts and do not value the company.

1. **Business model / positioning:** synthesize products, qualification, customers/end markets, direct/distributor route to market and geography. State only what the profile supports.
2. **Historical trajectory:** recompute revenue growth where possible. Treat volume and price/mix as disclosed contribution points to revenue growth. Interpret utilization and reported EBITDA margin without asserting an unproven causal bridge.
3. **Key tensions:** identify the material evidence gaps or management statements that require qualification.
4. **Three ranked questions + why each matters:** give exactly three questions. For each, name the later judgement it could change: revenue, margin, normalization, risk, peer selection or valuation.

For each material number or management statement, use a brief profile/table locator. Keep observed data, management statements and analyst inferences distinct.

Source basis: INSTRUCTOR SYNTHESIS: partial feedback after baseline; the P0–P4 slides explain the design changes later.

[Reference](#)

Four observable dimensions for the VP note

Dimension	Score 2: observable standard
Business-model synthesis	Connects product specification/qualification, customer/end-market mix, direct/distributor route to market and geography/footprint into a coherent operating configuration. It does not invent pricing power.
Historical economics	Reconciles revenue growth; distinguishes volume contribution from price/mix contribution; interprets utilization and reported EBITDA-margin movement; identifies association rather than unproven causality. Numerical fidelity belongs here.
Evidence and uncertainty discipline	Keeps observed data and management statements distinct; qualifies causal interpretation; imports no external facts; turns unsupported statements into questions where appropriate.
Prioritization / VP usefulness	Uses exactly three ranked questions. Each says which later judgement could change. The note is concise and decision-useful.

Common anchors: **0** absent, incorrect or unsupported; **1** present but partial, vague or mechanically copied; **2** meets the stated standard accurately and concisely.

Source basis: INSTRUCTOR SYNTHESIS: common rubric for baseline, redo and optional homework.

Reference

Partial feedback: an analytical route, not the completed answer

Route

1. Build a synthesis of the operating model.
2. Explain the historical change in drivers, rather than list annual snapshots.
3. Name tensions created by the supplied evidence.
4. Ask only questions whose answers could change a later valuation judgement.

Not shown yet

- complete model answer;
- complete ranked-question set;
- full fact/management/inference treatment;
- multi-turn instructor workflow;
- final source reconciliation.

Professional destination at the halfway point

Make a small evidence pack decision-useful while preserving what it does not establish.

Source basis: INSTRUCTOR SYNTHESIS: partial-feedback boundary and TargetCo source scope.

Reference

Worked interpretation: the revenue bridge is a bridge of contributions

	2022A	2023A	2024A	2025A
Revenue growth	–	2.1%	1.6%	5.9%
Volume contribution (ppt)	–2.5%	–4.0%	0.5%	4.0%
Price / mix contribution (ppt)	10.0%	6.1%	1.1%	1.9%
Utilization	82%	78%	74%	84%
Reported EBITDA margin	14.8%	14.0%	12.8%	14.3%

How to read the table. Positive price/mix contribution can offset falling volume. Recovering utilization can coexist with incomplete margin recovery. The table therefore shows observed co-movement, not a causal margin bridge.

A defensible reading

In 2023A, +6.1 ppt price/mix offset a –4.0 ppt volume contribution. In 2025A, +4.0 ppt volume is the larger disclosed contributor. Utilization recovered to 84%, but the 14.3% margin remains 50 bp below 2022A. Recovery is evident; the supplied pack does not quantify the margin bridge or prove a steady-state outcome.

Source basis: TARGETCO APPLICATION: supplied historical table and period commentary.

Reference

Partial feedback: business-model synthesis

What the pack supports

TargetCo is a specification-led manufacturer rather than merely a reseller. Engineering, tooling and customer qualification, combined with a 71% direct/key-account mix, indicate a service- and qualification-intensive route to market. The distributor channel and European footprint add commercial context.

What the pack does not establish

It does not establish pricing power, retention, channel margin or contract durability. The 9.4% largest-customer and 41.8% top-ten figures are facts requiring context, not an automatic risk classification.

Valuation-relevant question

What is the end-market, channel and facility-level evidence behind the 2025A volume recovery?

Its answer could change revenue confidence, risk and later peer selection.

Source basis: TARGETCO APPLICATION; INSTRUCTOR SYNTHESIS.

Reference

LLMs generate conditional text; they do not supply a truth test

Term	Working definition for this course
AI	Computational systems that perform tasks associated with perception, prediction, generation or decision support.
Generative AI	AI that produces new content, such as text, code, images or audio.
LLM	A generative model trained to predict and generate sequences of language tokens.
Token	A unit of text processed by the model; a token is not reliably a word.
Context window	The finite token sequence the model can use at one time: instructions, documents, prior messages and current task instruction.
Inference / generation	Running the already-trained model: given current context and tokens already generated, the model computes a conditional distribution over the next token; the decoding rule selects or samples a token and the process repeats until completion.

Analytical implication

An answer can be grammatically convincing while misstating a source, skipping a number or inventing a causal story. Language quality is not source fidelity.

Source basis: SOURCE-DERIVED: official OpenAI prompt guidance; Anthropic AI Capabilities and Limitations; INSTRUCTOR SYNTHESIS for course definitions.

Reference

Autoregressive generation explains why context changes outputs

A minimal technical model

$$p_{\theta}(x_1, \dots, x_T \mid c) = \prod_{t=1}^T p_{\theta}(x_t \mid x_{<t}, c).$$

- x_t is the next token; $x_{<t}$ is the preceding generated sequence.
- c is supplied context: task, source material, examples and prior conversation.
- θ denotes learned model parameters. At generation time, the application repeatedly chooses or samples a next token from a conditional distribution.

Consequence

Changing a prompt or context changes the distribution of possible outputs. It does not create a criterion for whether an output is true, financially reasonable or properly sourced.

Source basis: INSTRUCTOR SYNTHESIS: standard autoregressive language-model formulation; no transformer mathematics is assumed.

Reference

Context design: evidence, authority and salience

Three terms before the trade-off

Evidence: information from the permitted source set used to support a claim or calculation.

Source authority: the status assigned by the analyst/task specifying which source governs a claim when sources differ.

Salience: how strongly a supplied item effectively influences generation. Analyst importance does not guarantee model salience.

TargetCo contrast

The profile is **authoritative for this exercise**. M-01 and M-02 inside it remain **management statements**, not independently established facts.

More context can add evidence but also irrelevant instructions, competing sources, weak salience and verification cost. Identify the permitted source, preserve its boundary, then request extraction before synthesis.

Source basis: SOURCE-DERIVED: Anthropic Prompting Best Practices, long-context guidance; TARGETCO APPLICATION. [Reference](#)

Model, application and tool: a TargetCo pathway

TargetCo PDF → application handles/parses file and manages chat/context → underlying model generates response → calculator, code or file tool may retrieve, recompute or check → analyst reviews result

Object	Meaning in this pathway
Model	The underlying generative model that produces conditional token predictions. A displayed version can matter, but it is not the entire workflow.
Application	The interface/orchestration layer that may handle files, memory, system instructions, context and model routing. Comparing applications does not isolate only their models.
Tool	An external capability, such as a calculator, code execution or file operation. It can retrieve or recompute something, creating a separate source and checking obligation.

Source basis: SOURCE-DERIVED: Anthropic AI Capabilities and Limitations; INSTRUCTOR SYNTHESIS.

Reference

AI fluency assigns the human work before, during and after AI use

Anthropic's 4D framework

Delegation – decide whether, when and what to hand to AI.

Description – communicate goal, context and standards.

Discernment – assess quality and usefulness.

Diligence – take responsibility for use and final delivery.

Modes

Automation: AI performs a routine task.

Augmentation: AI expands a person's analysis while the person directs and judges it.

Agency: AI pursues a multi-step objective with delegated action scope.

Course map: Session 1 is principally **augmentation**: delegate extraction and drafting, retain judgement and verification. Session 6 moves closer to agency, with plans, approval points and stopping rules.

Source basis: SOURCE-DERIVED: Anthropic AI Fluency: Framework & Foundations; INSTRUCTOR COURSE MAP.

[Reference](#)

Success criteria and evaluation

Definitions

Success criterion: property the output must satisfy. **Evaluation:** observable test. **Failure example:** result showing the criterion was not met.

Criterion	Evaluation	Failure example
Historical economics	Check that 2023A says +6.1 ppt price/mix offsets −4.0 ppt volume and reconciles the displayed 2.1% growth.	Calls 2023A volume-led growth.
Evidence discipline	Inspect the treatment of M-02.	States commissioning costs are non-recurring as an established fact.
Business-model synthesis	Check whether qualification, route to market and footprint form an operating configuration.	Calls TargetCo a reseller or asserts pricing power.
Prioritization / VP usefulness	Count exactly three ranked questions and inspect the named later judgement.	Gives an unranked generic diligence list.

Source basis: SOURCE-DERIVED: Anthropic Define Success Criteria and Build Evaluations; TARGETCO APPLICATION.

[Reference](#)

Prompt ablation: P0 through P4

Version	Increment	What it is intended to change
P0	Student-composed baseline prompt from common sketch.	Establish the baseline workflow without an advanced method.
P1	Strong instructor single-shot instruction, shown after baseline.	Makes the permitted source and four output sections explicit; requests exactly three ranked questions.
P2	Analysis before drafting.	Makes extraction, contribution arithmetic and inference reviewable before prose hides an error.
P3	Explicit criteria and checks.	Uses the four observable success criteria on this slide and checks material source treatment.
P4	Extract, challenge, draft and verify across a multi-turn workflow.	Localizes errors and creates human handoff points.

Method

If a response uses information outside the permitted source or answers a different task, improve the specification and source boundary. Missing evidence or an unsuitable tool cannot be fixed merely by longer wording.

Source basis: SOURCE-DERIVED: Anthropic Prompt Engineering Overview; OpenAI Prompt Engineering Best Practices; INSTRUCTOR SYNTHESIS.

[Reference](#)

P1 and P2: specification before polished prose

P1: strong single-shot instruction

Use only the supplied TargetCo profile. Produce four sections: business model / positioning; historical trajectory; key tensions; three ranked questions + why each matters. No external facts or valuation.

What P1 changes

It makes source authority and the requested output structure visible. It still leaves the reasoning process hidden.

P2: analysis before draft

Before drafting, extract business-model evidence; recompute historical growth; distinguish volume and price/mix contribution points; track utilization and margin; list what each point supports and does not establish. No final prose yet.

What P2 changes

It creates a reviewable intermediate artifact. It does not make the model's interpretation independently true.

Source basis: SOURCE-DERIVED: Anthropic Prompting Best Practices, clear/direct instructions and sequential tasks; OpenAI Prompt Engineering Best Practices.

[Reference](#)

P3 and P4: an explicit test and a reviewable chain

P3: attach criteria and checks

Check the four observable success criteria: operating configuration; historical economics including contribution arithmetic; evidence/uncertainty discipline; exactly three ranked, valuation-relevant questions in a concise VP note.

Why chaining is a design choice

Advantages: intermediate outputs are reviewable; errors are easier to localize; humans can intervene.

Costs: more interactions and latency; state/context management; new opportunities for inconsistency between steps.

P4: make a chain

Extract evidence → challenge against criteria → draft VP note → reconcile to source.

Source basis: SOURCE-DERIVED: Anthropic Prompt Engineering Overview and Prompting Best Practices; INSTRUCTOR SYNTHESIS.

Reference

Examples, structure and roles

Technique	What it changes and when it helps	What it cannot do
Examples	A non-TargetCo output-shape example can demonstrate restrained phrasing or a compact question format when the requested form is ambiguous.	It cannot supply TargetCo evidence. A filled TargetCo example before baseline leaks the analytical answer.
Structured prompt	Clear headings separate task, sources, output and checks when a prompt has several components.	It cannot validate numbers or make a management statement true.
Role	"Write for a VP from the supplied profile" can set audience, tone and evidence discipline.	It cannot provide privileged facts, prove causality or transfer accountability.

Concrete role comparison

Weak: "You are the world's best banker. Tell me what will happen." **Useful:** "Write for a VP from the supplied material; label uncertainty rather than invent detail."

Source basis: SOURCE-DERIVED: Anthropic Prompting Best Practices, examples/XML/roles; OpenAI Prompt Engineering Best Practices; INSTRUCTOR SYNTHESIS.

Reference

Source architecture and structured prompts

Definitions

Task instruction = the question or work order the model must answer after receiving documents.

Source-grounded extraction = return the relevant fact, statement or calculation with its source location before asking the model to synthesize across documents.

Neutral portable form

Task: prepare a one-page TargetCo VP note.

Sources: supplied TargetCo profile only.

Output: four named sections.

Criteria: four observable success criteria.

Checks: recompute revenue bridge; reconcile material claims.

Both express the same logical structure. XML is syntax that can separate components in a complex Claude prompt, not a different reasoning method.

Anthropic XML form

```
<task>prepare a one-page TargetCo VP note </task>
```

```
<sources>supplied TargetCo profile only </sources>
```

```
<output>four named sections </output>
```

```
<criteria>four observable success criteria </criteria>
```

```
<checks>recompute bridge; reconcile claims </checks>
```

Source basis: SOURCE-DERIVED: Anthropic Prompting Best Practices, XML structure and long-context guidance.

[Reference](#)

Verification hierarchy

Level	Question	TargetCo example
1. Output / completeness	Did the draft use the four sections and exactly three questions?	Count sections and questions; check one-page form.
2. Source fidelity	Does every material number and statement reconcile to supplied evidence?	Confirm 2025A revenue = EUR 481.0m; preserve M-02 as management's statement.
3. Independent check	Can a calculation, source or tool not dependent on the generated prose verify the result?	Recompute $481.0/454.0 - 1 = 5.9\%$; inspect the source table or a separate spreadsheet.

Boundary

Asking the same model “are you sure?” may detect an omission. It is not Level 3 independent verification.

Source basis: INSTRUCTOR SYNTHESIS, aligned with the course Verification Framework; TARGETCO APPLICATION.

Reference

How AI can help make the distinction operational

Observed data	Management statement	Analyst inference
Supplied operating data or source content. Example: 2025A utilization = 84%.	What management says; it is informative but not independently established. Example: M-02 calls two years of costs non-recurring.	A qualified interpretation drawn from facts and statements. Example: recovery is evident, but the pack does not prove a steady-state margin.

Operational workflow

1. AI proposes a classification for each material statement. 2. AI supplies a source locator and rationale. 3. Analyst verifies source and decides whether inference is warranted. 4. Correct misclassified or unsupported claims before drafting.

This taxonomy is course instructor synthesis, not a Rosenbaum/Pearl definition.

Source basis: INSTRUCTOR SYNTHESIS; TARGETCO APPLICATION.

Reference

An evidence matrix records support and limits

Claim / source location	Observed data, management statement or calculation	Supports	Does not establish / open question
2025A growth character Historical table	Volume contribution +4.0 ppt; price/mix contribution +1.9 ppt; revenue +5.9%	2025A differs from 2023A	Breadth by customer, channel or end market
Utilization/margin tension 2022A and 2025A	Utilization 82% → 84%; margin 14.8% → 14.3%	Utilization alone is insufficient	Mix, cost, pricing and facility bridge
M-02 needs challenge Management material M-02	Management calls costs in 2024A and 2025A “non-recurring”	Recurrence is valuation-relevant	Whether any cost is recurring or a justified adjustment

The matrix is an analytical workpaper: claim, source, classification/calculation, support and limit. It makes uncertainty explicit rather than hiding it in polished prose.

Source basis: INSTRUCTOR SYNTHESIS; TARGETCO APPLICATION.

Reference

Redo: a new chat, the same source pack

15-minute redo

Start a **different fresh ephemeral chat**. ChatGPT groups use a new non-personalized Temporary Chat where available. Claude groups use a new Incognito chat. AI groups remain with their application; former human-only groups receive AI access and may split between applications.

Maximum three user messages

1. evidence / analysis;
2. challenge + draft;
3. verify + revise.

Do not send a message to erase context. The fresh chat separates redo from baseline.

Record

- actual baseline prompt;
- fresh-chat mode and three prompts/labels;
- delegated versus retained judgement;
- challenged statement; numerical/source check; remaining uncertainty;
- short baseline-versus-redo reflection.

Source basis: INSTRUCTOR SYNTHESIS; official OpenAI/Anthropic product locators in Session 1 source notes.

Reference

Post-redo answer I: synthesis and historical economics

Business model / positioning

TargetCo is a European specialty manufacturer of technically specified rigid containers, closures and dispensing assemblies. Engineering, tooling and customer qualification form part of the offer. Direct/key accounts represent 71% of FY2025A revenue; selective distributors represent 29%. Food & Beverage is the largest end market (55%), followed by Healthcare & Pharmaceutical (24%) and Premium Consumer (21%); France (32%) and DACH (24%) are the largest disclosed geographies. The 9.4% largest-customer and 41.8% top-ten measures describe concentration only: the pack provides no benchmark, identities, contracts or retention evidence.

Historical trajectory

Revenue rose from EUR 438.0m in 2022A to EUR 481.0m in 2025A. In 2023A, +6.1 ppt price/mix offset a -4.0 ppt volume contribution; this differs from 2025A's +5.9% growth, where +4.0 ppt volume is the larger disclosed contributor. Utilization fell to 74% in 2024A during Łódź commissioning and recovered to 84% in 2025A; reported EBITDA margin recovered from 12.8% to 14.3%, still below 2022A's 14.8%. Recovery is supported, but the pack does not prove its full causal bridge.

Source basis: TARGETCO APPLICATION: full professional answer, Part I.

Reference

Post-redo answer II: tensions and ranked questions

Why these are tensions

- **Utilization versus margin:** 2025A utilization exceeds 2022A, but margin is 50 bp lower. If utilization is a main recovery lever, an unobserved margin bridge remains.
- **Management recovery bridge:** M-01 cites utilization, mix and efficiency, but the pack does not quantify their contributions.
- **“Non-recurring” across two years:** two-year duration weakens automatic one-off treatment, but does not prove recurrence.

Three ranked questions – and why

1. **What bridges margin from 2022A to 2025A and supports run-rate?** It could change margin forecasts and valuation.
2. **What is the nature and duration of 2024A–2025A commissioning/quality costs?** It could change normalization judgement.
3. **How broad and durable is volume recovery by customer, channel, end market and facility?** It could change revenue confidence, risk and peer selection.

Evidence limitation

M-01 and M-02 remain management statements. The supplied pack supports neither a forecast nor a valuation conclusion.

Source basis: TARGETCO APPLICATION: full professional answer, Part II.

Reference

Instructor workflow: AI helps when its work is challengeable

1. **Evidence / analysis:** work from the slide-29 tensions. Extract business-model evidence; recompute revenue growth; identify contribution points; track utilization and margin. Return source location, support and limit. No final prose.
2. **Challenge + draft:** test the analysis against four criteria, correct unsupported inference or generic questions, then draft the one-page note in four sections.
3. **Verify + revise:** reconcile every number and material statement to source or calculation. Revise only where the evidence supports revision; flag unresolved items.

The analyst still determines

Source authority; whether causal language is justified; which tension matters for valuation; whether a cost qualifies for normalization; whether the final answer is acceptable.

Source basis: INSTRUCTOR SYNTHESIS: full prompts in `instructor/SESSION1_INSTRUCTOR_AI_WORKFLOW.md`.

Reference

Compare the actual outputs

Baseline versus redo

- Which answer correctly decomposes the historical revenue story?
- Which recognizes the utilization/margin tension?
- Which question would change later valuation assumptions most?
- Which output reports a management assertion as if independently verified?
- Which output invents a causal explanation unsupported by the pack?

This remains an observed application comparison, not a timeless vendor ranking.

Optional homework: same-prompt

cross-application comparison Use the same instructor near-final prompt sequence, same TargetCo profile and fresh Temporary/Incognito chats. Match tools, score both on the common rubric, record date/application/displayed model-version/tool-memory state, then write what differences matter to a VP.

Source basis: INSTRUCTOR SYNTHESIS: `instructor/BASELINE_DISCUSSION_GUIDE.md`.

Reference

Each session changes the work product, not the accountability

Session	Professional work product
1	Target-study note.
2	Audited Excel workbook and change log.
3	Valuation analysis.
4	Reconciled Python/Excel analysis.
5	Checked deal-extraction table plus derived transaction indicators.
6	Strategic buy-side briefing/workflow.

Continuity

The finance artifact, technical environment and permitted AI role become more demanding. The analyst retains accountability for specification, evidence, judgement and verification.

Session 2 preparation

Next professional question

Can we rely on TargetCo's historical Excel model? The task moves from target-study narrative to financial-statement and model logic; it does not introduce valuation.

Refresh

- reported/adjusted EBITDA; debt, cash, net debt and leverage;
- CapEx/NWC consistency and formula tracing;
- Rosenbaum & Pearl, *Investment Banking*, Ch. 1: "Financial Profile" and "Calculation of Key Financial Statistics and Ratios."

Use these sections to refresh the concepts before class; you are not expected to reproduce or solve the textbook exercises unless explicitly assigned.

Use the verified locators

- Workbook Ch. 1: business/financial profile Q14–21; growth/profitability/credit Q7–9; one-time/non-recurring Q3.
- Koller, Goedhart & Wessels (2025), Chs. 11 and 13, for optional model organisation/hygiene context.

Source basis: SOURCE-DERIVED: verified private instructor-copy locators; no textbook exercise is reproduced.

Reference

The analyst owns the specification and the check

Today

A strong target study synthesizes business configuration, operating evidence and valuation-relevant uncertainty. AI can accelerate the path from materials to a draft; it cannot decide which sources govern or whether an interpretation is justified.

Evidence → analysis → challenge → draft → reconciliation

Before next session

Keep your group AI Work Record and final prompt-chain locator. In Session 2, the question becomes whether a historical Excel model can be trusted after inspection and independent checks.