

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:** 150-minute final in-class examination.
- **Office hours:** feel free to (DO) come: B612.
- Slides/materials are updated regularly. Report a typo/mistake: corrections help the class.
- **This session:** a first-pass sell-side target study before valuation, with AI as augmentation.
- **Key references:** i) Rosenbaum & Pearl, Chapter 1 “Study the Target” (Exhibit 1.3); ii) official OpenAI and Anthropic guidance; iii) Anthropic AI Fluency.

Learning objectives

During your internship...

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.

Today's learning sequence

1. Study TargetCo without a prompting method.
2. Compare outputs and (partial) feedback.
3. LLM and disciplined workflow.
4. Redo the task in a fresh chat, then inspect the full answer and workflow.

Structure of our course: 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 choices.

Business profile

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

Financial profile (selected dimensions)

- size;
- profitability;
- growth;
- (returns and credit profile later).

Source basis: Rosenbaum & Pearl (2013), Ch. 1, Step I “Study the Target” (Exhibit 1.3).

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 vs. margin: they can move differently, as margin also depends on cost recovery, utilization, mix and other operating factors.

TargetCo: 2025A utilization exceeds 2022A, while the 2025A reported margin is lower.

↪ analytical tension (not clear causality).

Baseline design: today's first assignment

	ChatGPT group	Claude group	Human-only group
Inputs	Same two-page TargetCo profile; same VP assignment; same time.		
Baseline	Student-composed 1 message prompt from common sketch.	Same finance assignment.	Same finance assignment w/out AI.
Redo	Temporary chat (max 3 messages).	Incognito chat (max 3 messages).	Possible AI access.
Main comparison	Baseline output/work process versus redo output/work process		

Interpretation discipline

The baseline is not a ChatGPT vs. Claude comparison: students may formulate different prompts.
↪ 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 basic starting sketch before sending your one 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: NO second message allowed.

20-min 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.

Start now. We compare outputs in about 20'.

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: Spoiler of the full sequence (three messages).

Reference

Four observable dimensions for the VP note

Dimension	Score 2/2:
Business-model synthesis	Connects product specification, 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.
Certainty vs. Uncert.	Keeps observed data and management statements distinct; qualifies causal interpretation; imports no external facts; turns unsupported statements into questions where appropriate.
VP usefulness	Uses exactly three ranked questions. Each says which later judgement could change. The note is concise and decision-useful.

Rate yourself: **0** absent, incorrect; **1** present but partial, or mechanically copied; **2** meets standards.

Source basis: Common rubric for baseline, redo and optional homework.

[Reference](#)

Partial answer: an analytical route

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 doc.
4. Ask only questions whose answers could change a later valuation judgement.

Yet to come:

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

The key for a good starting point:

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

Working interpretation about revenues

	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 co-movement, not a causal margin!

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: Full 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 it does not support/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: What the two-page Target Co profile can support

Reference

LLMs generate conditional text, not truth

Term	Definition
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 (e.g., check "AI allucination"). Language quality is not source fidelity.

Source basis: Model, Application Tool in workflow

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.

Context design: evidence, authority and salience

A bit of terminology:

Evidence: information from the (restricted) source used to support a claim/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 imply 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: Long-context failures

Reference

Model vs. Application vs. Tool

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 (e.g., Codex 5.6 Terra).
Application	The interface/orchestration layer that may handle files, memory, system instructions, context and model routing (e.g., Claude code vs. Claude browser chat).
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.

Source basis: Model, Application Tool in workflow

Reference

AI fluency = human work before, during and after AI use

Anthropic's 4D framework

Delegation – decide whether/when/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: Session 1 and AI fluency 4Ds

Reference

Success criteria and evaluation

Being clear about criteria:

Success criterion: property the output must satisfy VS. **Evaluation:** observable test VS. **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: Evaluation comes before prompt cleverness.

Reference

Prompt trajectory: P0 through P4

Prompt	Increment	What it is intended to change
P0	Student-composed baseline prompt from common sketch.	Establish the baseline workflow without an advanced method.
(P1)	(Instructor's single-shot instruction.)	(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: Prompt-chain skeleton.

[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: Prompt-chain skeleton.

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; paid vs. unpaid versions.

P4: make a chain

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

Source basis: Prompt-chain skeleton.

[Reference](#)

Missunderstanding? Examples, structure and roles

Technique	What it changes and when it helps	What it cannot do
Examples	When the form is ambiguous: a non-TargetCo output-shape example can demonstrate phrasing or a compact question format.	It cannot (and should not) 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: Examples of Appropriate vs. Inappropriate.

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.

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: Example of neutral vs. anthropic.

[Reference](#)

Verification hierarchy

Level	Question	TargetCo example
1. 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 an independent verification.

Source basis: Self-check vs. independent check.

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 reference + rationale. 3. Analyst verifies source and decides about the inference. 4. Correct misclassified or unsupported claims before drafting.		

An evidence matrix records evidence and especially 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 calls costs in 2024A and 2025A “non-recurring”	Recurrence is valuation-relevant	Whether any cost is recurring or a justified adjustment

The matrix dissects: claim, source, classification/calculation, support and limit.

It makes uncertainty explicit rather than hiding it in polished prose.

Source basis: Full(er) evidence matrix.

Reference

Redo: a new chat, the same source pack

15-minute redo

Start a **different fresh chat**.

ChatGPT groups use Temporary chat (where available).

Claude groups use Incognito chat (where available).

Former human-only groups receive AI access and may split between applications.

Maximum 3 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.

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: Full(er) evidence matrix.

Reference

Post-redo answer II: tensions and ranked questions

Why these are tensions

- **Utilization vs. Margin:** 2025A utilization exceeds 2022A, but margin is 50 bp lower. If utilization is a main recovery lever, an unobserved margin bridge remains.
- **(Management-inferred) Recovery:** 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

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: Why the three questions are ranked.

Reference

Instructor workflow: AI helps when its work is challengeable

1. **Evidence / analysis:** work from the previous slide's 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's prompt sequence. Full prompt in `SESSION1_INSTRUCTOR_AI_WORKFLOW.md`.

Reference

Take some time to 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?

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.

Each of our sessions changes the work product: not 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 (not yet 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.

Book references

- 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.

Appendix: target-study profile and later valuation work

Business profile: why each dimension matters

- Sector/subsector frames drivers, cyclicity, etc.
- Products/services distinguish commodity-like supply from specified/service-intensive offers.
- Customers/end markets distinguish purchaser from demand exposure.
- Channels/geography can alter commercial model, cost base and comparability.

Financial profile: why selected dimensions matter

- Size can correlate with scale and market comparability.
- Profitability describes conversion of sales into profit.
- Growth needs a driver-level reading.
- Returns/credit require later calculations.

Source basis: Rosenbaum & Pearl (2013), Ch. 1, Step I and Exhibit 1.3; original paraphrase.

[Back to core](#)

Appendix: what the two-page TargetCo profile can support

Evidence supplied	A cautious analyst may say	A cautious analyst must stop short of saying
End-market, geography, channel and concentration percentages	Reported mix and configuration follow-up.	Concentration is high/low, commercially good/bad or contractually secure.
Revenue bridge, utilization and margin history	2023A price/mix contribution offsets falling volume; 2025A has a larger volume contribution; utilization and margin recover from 2024A.	A quantified causal connection, or sustainable run-rate.
M-01 and M-02 management statements	What management says about recovery levers and cost classification.	That recovery will occur or an adjustment is justified.

Classroom purpose

The analyst must synthesize, calculate, challenge and prioritize.
The short deck is deliberately incomplete!

Source basis: Back.

[Back to core](#)

Appendix: four-dimensional rubric calibration

Dimension	0	1	2
Business-model synthesis	Copies labels or omits configuration.	Names facts but weakly connects them.	Connects specification/qualification, customer/end market, route to market and footprint without inventing economics.
Historical economics	Wrong contribution story, arithmetic or causality.	Lists figures with a weak bridge.	Reconciles growth, distinguishes contributions, treats utilization/margin cautiously and keeps numerical fidelity.
Evidence and uncertainty discipline	Imports facts or treats management assertion as fact.	Notes uncertainty but blurs categories.	Separates observed data, management statements and qualified inference; turns unsupported claims into questions.
Prioritization / VP usefulness	Not three questions or generic list.	Three questions but weak consequence or unclear note.	Exactly three ranked questions; each names a later judgement; concise VP-ready note.

Source basis: Back.

[Back to core](#)

Appendix: model, application and tool in the TargetCo workflow

Concrete pathway

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

Term	Operational consequence
Model	Produces conditional next-token predictions. Record any displayed model/version; do not infer hidden configuration.
Application	May add file conversion, system instructions, context management, routing, memory or tool access. An application comparison is not a pure model experiment.
Tool	Retrieves, calculates or changes an external object. Tool output creates a source and checking obligation.

Source basis: Back.

[Back to core](#)

Appendix: long-context failure modes

Failure mode

- management material and an authoritative source are treated as equally established;
- source material contains instructions that compete with the task;
- material evidence is lost among irrelevant pages;
- a citation label appears without a reconciled claim.

Design response

- label source, date and authority;
- delimit documents and permitted use;
- make the task instruction easy to locate;
- extract fact/statement/calculation with location before synthesis.

TargetCo lesson

The profile is authoritative for the exercise while M-01/M-02 within it remain management statements. Authority and truth status are different analytical questions.

Source basis: Back.

[Back to core](#)

Appendix: evaluation comes before prompt cleverness

Evaluation loop

task-specific test cases → initial prompt → run → score → diagnose → revise.

Observed failure	Precise diagnosis	Response
Calls 2023A demand/volume-led	The output ignores disclosed −4.0 ppt volume and +6.1 ppt price/mix contributions.	Require source-grounded extraction and contribution recomputation before prose.
Treats M-02 as fact	Management-statement classification/source-fidelity check is missing.	Request classification, locator and analyst verification before drafting.
Gives six generic questions	Prioritization criterion is not operational.	Require exactly three ranked questions and one named valuation consequence each.
Correct but too long	Output constraint fails.	State one-page constraint and use a non-TargetCo output-shape example if format remains unstable.

Source basis: Back.

[Back to core](#)

Appendix: few-shot examples and the baseline boundary

Strong non-TargetCo output-shape example

Observed fact: a product requires customer qualification before serial production.

Restrained interpretation: qualification may make switching and implementation relevant to diligence.

Open question: what evidence shows the duration, economics and renewal pattern of qualified relationships?

Appropriate Use this type of example after baseline when a group needs a model of fact → qualified interpretation → question.

Inappropriate before baseline A filled TargetCo example leaks the analytical answer, anchors wording and turns the task into transcription.

Source basis: Back.

[Back to core](#)

Appendix: the same semantic prompt in headings and XML

Neutral headings

Task

Prepare a one-page TargetCo VP note.

Sources

Supplied TargetCo profile only.

Output

Four named sections.

Criteria

Four observable success criteria.

Checks

Recompute bridge; reconcile claims.

The tags separate task, source boundary, requested output, criteria and checks. They do not create a different reasoning method or independently verify a result.

Anthropic XML

```
<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: Back.

[Back to core](#)

Appendix: roles set audience and standards

Prompt A

You are a senior investment banker. Explain TargetCo's prospects and give a valuation view.

Problem: the role conflicts with the assignment, invites external knowledge and encourages an unsupported valuation claim.

Prompt B

Write for a VP considering a sell-side mandate. Use only the supplied profile. State the operating picture, label management statements and leave unsupported causality as questions. Do not value the company.

Benefit: audience, vocabulary, source boundary and standards are explicit.

Rule

A role guides tone, context and perspective. It cannot add facts, make an inference true or transfer accountability.

Source basis: Back.

[Back to core](#)

Appendix: prompt-chain skeleton

Stage	Output	Human question	Failure localized
Extract	Evidence table; calculations	Right facts, contribution points and source locations?	Retrieval / arithmetic
Challenge	Criteria-based corrections	What is unsupported, omitted or generic?	Interpretation / specification
Draft	One-page VP note	Is it concise, qualified and decision-useful?	Communication / prioritization
Verify	Reconciliation log	Does each material point return to a source or calculation?	Source fidelity

State-management cost

Each stage needs a stable source pack, versioned prompt and handoff. The live redo compresses the four logical stages into at most three user messages.

Source basis: Back.

[Back to core](#)

Appendix: self-check versus independent check

Check request	What it can reasonably do	What is still required
"Review this draft against the rubric."	Detect missing sections, generic questions or an obvious internal contradiction.	Reconcile each material statement to the supplied profile.
"Check the calculations."	Recalculate a percentage or identify arithmetic inconsistency.	Independently recompute in calculator/spreadsheet and identify inputs/source.
"Are you sure M-02 is non-recurring?"	Surface qualification language or admit missing support.	Obtain evidence on cost nature/duration; analyst decides later normalization treatment.

TargetCo implication

M-02 spans two years. That supports a question about recurrence and duration; it neither proves recurrence nor non-recurrence, and it establishes no adjustment amount.

Source basis: Back.

[Back to core](#)

Appendix: fuller TargetCo evidence matrix

Analytical point	Source location	Evidence/calculation	Support, limit and next question
Specification-led configuration	Company extract; product/facility notes	Qualification before serial production; containers, closures and dispensing components; 71% direct/key account mix	Supports specialty manufacturing/service configuration; does not show pricing power or channel margin.
2023A growth character	Historical table; 2023A commentary	Volume contribution −4.0 ppt; price/mix +6.1 ppt; revenue +2.1%	Supports contrast with 2025A; no customer-level price/mix bridge.
2025A utilization/margin tension	2022A–2025A table	Utilization 82% to 84%; margin 14.8% to 14.3%	Supports challenge to single-driver causality; open cost/mix/facility bridge.
Customer concentration	Commercial snapshot	Largest 9.4%; top ten 41.8%	Supports request for contract and retention context; no automatic risk conclusion.
M-01/M-02	Management materials	Asserted recovery levers; two years of stated non-recurring costs	Supports diligence questions; does not establish forecast or adjustment.

Source basis: Back.

[Back to core](#)

Appendix: why the three questions are ranked

1. **Margin bridge and supported run-rate.** The utilization/margin tension leaves an unobserved bridge. The answer could alter margin forecasts and valuation.
2. **Nature and duration of commissioning/quality costs.** Identifiable, finite, incremental and non-run-rate evidence affects whether any normalization is supportable. Two-year duration creates a challenge; it does not establish a treatment.
3. **Breadth and durability of volume recovery.** Evidence across customer, channel, end market and facility can change forecast confidence, risk and peer selection.

Acceptable alternative

A question about qualification timing or direct-account/distributor economics can be equally strong if it states a potential valuation consequence and is prioritized against these questions.

Source basis: Back.

[Back to core](#)

Appendix: instructor prompt-chain skeleton

Message 1: evidence / analysis

Use only the supplied TargetCo material. Extract business-model evidence; recompute revenue growth; distinguish disclosed volume and price/mix contribution points; track utilization and margin; identify tensions; return source location, support and limit. No valuation or final prose.

Message 2: challenge + draft

Check the analysis against the four criteria. Correct unsupported inference, omitted dimensions, numerical inconsistency and generic questions. Then draft the one-page note under Business model; Historical trajectory; Key tensions; Three ranked questions + why.

Message 3: verify + revise

Reconcile each material number and statement to source or calculation. Revise the note only for verified corrections. Flag remaining uncertainty rather than supplying a fact.

Source basis: Back.

[Back to core](#)

Appendix: a strong near-final prompt sequence for TargetCo

Message 1 – evidence / analysis. The supplied TargetCo profile is authoritative. Use no external facts, tools or valuation. Extract product/qualification, markets, channel and footprint; revenue and $g_t = R_t/R_{t-1} - 1$; disclosed volume and price/mix contribution points (check $g_t = v_t + p_t$); utilization and reported margin; and M-01/M-02 as management statements. For each, give profile location, classification, calculation/support and limit. Flag 2025A's 84% versus 82% utilization but 14.3% versus 14.8% margin tension. Do not draft.

Message 2 – challenge + draft. Assess the extraction against four criteria: coherent business-model synthesis; correct contribution arithmetic and qualified causality; evidence/uncertainty discipline; concise VP usefulness with exactly three ranked questions naming a later judgement. Correct only what the profile supports. Draft one page: Business model / positioning; Historical trajectory; Key tensions; Three ranked questions + why. Preserve M-01/M-02 as management statements. No valuation.

Message 3 – verify + revise. Reconcile every material number and statement to a profile location or calculation. Recheck 2023A ($-4.0 \text{ ppt} + 6.1 \text{ ppt} = 2.1\%$) and 2025A ($4.0 \text{ ppt} + 1.9 \text{ ppt} = 5.9\%$). Confirm that no management statement became fact, causal language remains qualified and exactly three questions are ranked. Give a short reconciliation list, revise only where required, and flag unresolved uncertainty.

Source basis: Back.

[Back to core](#)

Appendix: common weak moves are diagnostic

Weak move	Why it fails and how to repair it
Rewrites profile headings	It has not synthesized an operating configuration. Connect product, qualification, channel and geography cautiously.
Lists annual data	It misses changing economics. Contrast 2023A price/mix contribution offsetting falling volume with 2025A's larger volume contribution.
Says utilization "caused" margin recovery	The pack has no quantified causal bridge. Use qualified language and ask for the bridge.
Calls M-01/M-02 facts	It converts management language into established evidence. Label the statement and specify missing support.
Calls 9.4% / 41.8% high or low	Benchmark and contractual context are absent. State the figures and ask for context.
Asks generic diligence questions	A VP cannot see valuation relevance. Name the later judgement that might change.

Source basis: Back.

[Back to core](#)

Appendix: Session 1 and the AI Fluency 4Ds

4D	Session 1 action	Evidence retained
Delegation	Delegate extraction, cross-check suggestions and drafting; retain authority, valuation relevance and approval.	Work Record: delegated versus retained.
Description	State task, authoritative source, output, criteria and checks.	P0 and P1–P4 versions.
Discernment	Score four dimensions; identify unsupported stories and omissions.	Rubric, challenge, baseline/redo.
Diligence	Reconcile sources; disclose uncertainty; do not treat management language as verified.	Matrix, numerical check, remaining uncertainty.

Mode

This is augmentation: AI extends analysis under human direction. Later agency needs scope, approvals, state and stopping rules.

Source basis: Back.

[Back to core](#)

Appendix: finance locators and boundaries

Rosenbaum & Pearl (2013), *Investment Banking*

Ch. 1, Step I “Study the Target”; Exhibit 1.3 Business and Financial Profile Framework; “Financial Profile”; and “Calculation of Key Financial Statistics and Ratios.” Used for target-study sequence, profile dimensions and Session-2 navigation.

Rosenbaum & Pearl Workbook (2013) and Koller et al. (2025)

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

Rights boundary

Slides paraphrase methods and cite locators. They do not reproduce protected wording, tables, exhibits, cases, exercises, answers, layouts or values.

Source basis: Back.

[Back to core](#)

Appendix: sources, claims and application

Label	Meaning in this deck
SOURCE-DERIVED	A principle is paraphrased from an identified official page or a textbook section/exhibit locator.
INSTRUCTOR SYNTHESIS	A teaching framework, rubric, prompt chain, explanation or wording designed for this course. It is not presented as a source quotation.
TARGETCO APPLICATION	The source-derived or instructor-synthesis method is applied only to disclosed canonical TargetCo facts and labelled management statements.

Source map

SLIDE_SOURCE_MAP.md records exact locators and adopted principles by internal slide anchor.

SESSION1_SOURCE_NOTES.md records source access and rights boundaries. The map enables review; it does not replace source reading.

Source basis: Back.

[Back to core](#)