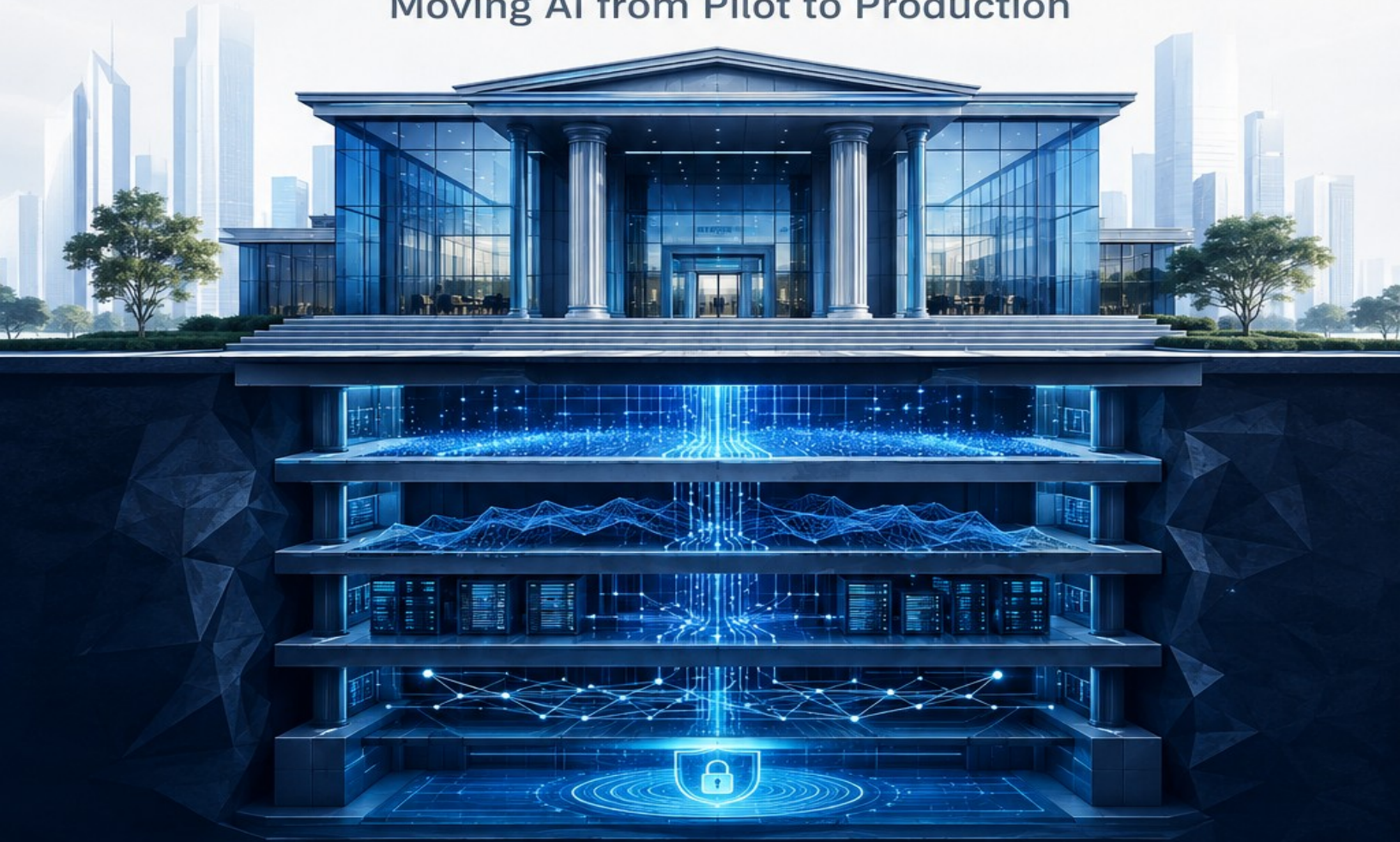


BUILDING THE AI-POWERED BANK

The Visual Engineering Playbook for
Moving AI from Pilot to Production



YOUR AI MODEL IS ONLY **20%** OF THE BANK.

ROMI NUR ISMANTO

BUILDING THE AI-POWERED BANK

The Visual Engineering Playbook for Moving AI from Pilot to
Production

Romi Nur Ismanto

AI Research Engineer · Jakarta

English Visual Edition · First Edition · 2026

Copyright © 2026 Romi Nur Ismanto. All rights reserved.

No part of this publication may be reproduced or distributed without written permission, except for brief quotations used in reviews or scholarly work.

This book is an independent educational and engineering reference. It is not legal, regulatory, investment, cybersecurity, or financial advice, and it is not an official publication of any vendor, bank, or regulator.

Regulatory and online references were reviewed through 12 August 2026. Readers should consult the latest official sources before implementation.

rominur.com · rominur@gmail.com

Preface

Banking has reached its most consequential technology transition since internet banking. Artificial intelligence is moving from innovation laboratories into the operating core: detecting fraud in milliseconds, assisting employees, personalizing customer journeys, interpreting documents and video, and supporting decisions that must remain explainable and auditable.

This book began with NVIDIA's public "Building the AI-Powered Bank" vision map, which depicts nine capabilities operating inside one fictional bank. The image is compelling because the use cases are visible. The harder engineering work is largely invisible: governed data, resilient integration, low-latency serving, MLOps, cybersecurity, model risk management, privacy, and human accountability.

The central proposition of this book is that the model is only a small part of a production banking capability. The 20/80 language used on the cover is a design heuristic, not an audited industry statistic. Its purpose is to force the right investment conversation: most durable value comes from the shared operating stack and the institution that can run it safely.

Each use-case chapter therefore connects four views: the business outcome, the model and data design, the end-to-end architecture, and the controls required in a regulated environment. The result is intended to serve executives who allocate capital, architects who design platforms, engineers who operate them, and risk professionals who must challenge and validate them.

Banking is data-rich, tightly regulated, and intolerant of silent error. A biased credit model, an unsafe agent, or an unexplained fraud decision can harm customers and the institution. Responsible AI is therefore treated as an engineering property of the entire system rather than as a final review step.

Use this book as a map, not a recipe. Adapt every pattern to your institution's products, customers, regulations, risk appetite, and operating model. Start with one measurable outcome, build the reusable foundation beneath it, and earn the right to scale.

About the Author

Romi Nur Ismanto

AI Research Engineer · Jakarta

rominur.com · rominur@gmail.com

Education

M.Eng. in Network Security

Universitas Indonesia, Jakarta · 2017

B.Eng. in Informatics

Universitas Gadjah Mada, Yogyakarta · 2004



How to Read This Book

This book is organized into four large parts which can be read sequentially or skipped as needed.

Part I — Foundations (Chapters 1–4) builds the foundation: why AI is a strategic imperative for banks, the core technologies behind modern AI (machine learning, deep learning, computer vision, NLP, generative AI), the computing infrastructure that underpins it (GPU, cloud, edge), and the data architecture that fuels it. Readers with experience in AI can skim through this section, but Chapter 4 on banking data is still recommended as it contains industry-specific context.

Part II — Nine Use Cases (Chapters 5–13) is the heart of the book. Each chapter dissects one capability from the "Building the AI-Powered Bank" infographic with a uniform structure: definition and business context, technical workings, reference architecture, data and model requirements, success metrics, implementation stages, risks and mitigations, as well as a practical checklist at the end of the chapter. This uniform structure is intentional so that chapters can be compared with each other and used as a template for internal feasibility studies.

Part III — Engineering and Operating (Chapters 14–18) brings it all together: the end-to-end architecture of an AI-powered bank, MLOps practices for managing dozens of models at once, security and responsible AI, the regulatory landscape (including the Indonesian context: PDP Law, POJK, and Bank Indonesia provisions), and implementation strategies from roadmap to ROI calculation.

Part IV — Lessons and the Agentic Future (Chapters 19–20) combines composite cases with publicly disclosed bank deployments, then develops the architecture, autonomy levels, operational controls, and human accountability required for agentic AI.

The back matter includes discussion questions, implementation templates, operational runbooks, a glossary, a chapter-by-chapter source map, a consolidated bibliography, and a subject index for professional reference.

Some writing conventions: technical terms that are standard in English are retained in their original form (e.g. feature store, fraud detection, inference) with explanations on the first occurrence; product and technology names (e.g. CUDA, TensorRT, Kafka) are written as is; and examples of code snippets or pseudocode are presented as necessary to clarify concepts, not as production code.

Who This Book Is For

This book is aimed primarily at technology practitioners in the financial industry and beyond:

- Solutions architects and enterprise architects who must design an end-to-end banking AI platform and integrate it with existing core banking.
- Data scientists and machine learning engineers who build and deploy fraud detection, recommendation, computer vision, and conversational AI models in highly regulated environments.
- Infrastructure, DevOps, and platform engineers setting up GPU clusters, data pipelines, and low-latency inference systems.
- Risk managers, IT auditors and compliance officers who need to understand how AI models work in order to monitor them effectively.
- Product leaders and digital transformation executives who must prioritize investments, build roadmaps, and remain accountable for results.
- Students and researchers in the fields of informatics, information systems and finance who want to see how AI theory is applied to real industry in all its complexity.

Reading prerequisites are intentionally light: a basic understanding of information systems and programming is sufficient. Machine learning concepts are explained from the ground up in Chapter 2, so readers without formal AI background can still follow the entire book.

Choose Your Reading Path

The book is designed for selective reading. Choose the route that matches the decision you need to make, then use the subject index to move across disciplines.

Reader	Start Here	What You Will Take Away
Executive / Board	Ch. 1, 14, 16–20	Investment logic, target architecture, risk ownership, roadmap, value realization, and the boundaries of agent autonomy.
Architect / Engineer	Ch. 2–15, 20	Data, GPU and edge infrastructure, serving, integration, MLOps, resilience, AgentOps, and production controls.
Risk / Compliance	Ch. 4, 9, 11–12, 16–20	Model risk, privacy, security, fairness, explainability, regulation, evidence, approvals, and auditability.
Use-Case Team	Ch. 5–13 + App. E	Nine reusable delivery blueprints with architecture, metrics, implementation stages, risks, and pilot checklists.

Contents

PART I — FOUNDATIONS		PART III — ENGINEER & OPERATE	
1 The AI-Powered Bank Era	8	14 End-to-End Architecture	79
2 Technology Foundations	14	15 MLOps	83
3 Accelerated Computing	23	16 Security and Responsible AI	87
4 Data Foundation	29	17 Regulation and Compliance	92
PART II — NINE USE CASES		18 Strategy and Execution	96
5 Personalized Branch Banking	35	PART IV — LESSONS & AGENTIC FUTURE	
6 Virtual Financial Advice	41	19 Cases and Patterns	103
7 Digital Twins	45	20 Agentic AI and the Future	108
8 Avatar Assisted Banking	49	REFERENCE MATERIAL	
9 Intelligent Video Analytics	54	Discussion Questions	113
10 Risk Visualization	59	Appendices A–G	121
11 Real-Time Fraud Detection	63	Glossary	147
12 Automated Investment Advisor	69	Source Map and Bibliography	154
13 ATM Security and Monitoring	74	Subject Index	159

Page numbers exclude the unnumbered front cover.

PART I — FOUNDATIONS

This section lays the foundation for the rest of the book: the strategic context for why banks around the world are racing to adopt AI, the core technologies that drive it, the computing infrastructure that underpins it, and the data that fuels it. The four chapters in this section form a common vocabulary that will be used continuously in the use case chapters.

VISUAL TAKEAWAY

The Bank Is a Governed System - Not a Model

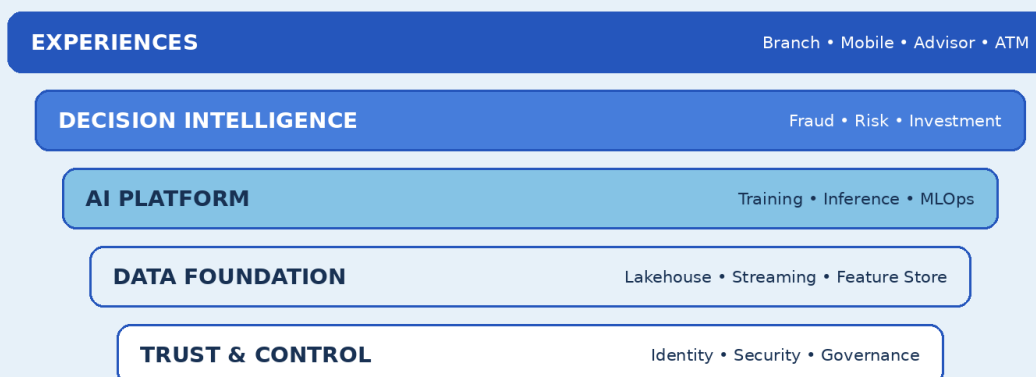


CHAPTER 1 — The AI-Powered Bank Era

FIGURE 01

The AI-Powered Bank Operating Stack

Shared foundations turn isolated models into an enterprise capability.



BUILDING THE AI-POWERED BANK • ENGLISH VISUAL EDITION

Figure 1. The AI-Powered Bank Operating Stack — Shared foundations turn isolated models into an enterprise capability.

VISUAL TAKEAWAY The strategic unit is not the model; it is the governed operating stack that repeatedly delivers safe decisions.

1.1 From Digitalization to Intelligence

Over the last three decades, banking has gone through three waves of technological transformation. The first wave was computerization: moving ledgers from paper to core banking systems, automating transaction recording, and connecting branches through private networks. The second wave is the digitalization of channels: internet banking, mobile banking, and open APIs that move customer interactions from the teller to the phone screen. The third wave — underway now — is the wave of intelligence: banks no longer simply record and channel transactions, but understand, predict and act proactively based on data.

The fundamental difference between the third wave lies in the nature of the system. Traditional digital systems are deterministic: the rules are written explicitly by the programmer ("if the balance is less than X, reject the transaction"). AI systems are probabilistic and learn from data: the model studies the patterns of millions of historical transactions to estimate the probability that a new transaction is fraudulent, without anyone writing down explicit rules. It is this shift from "written rules" to "learned patterns" that enables banks to handle complexities that would be impossible for

rules-based systems: ever-changing customer behavior, evolving modes of crime, and exponentially growing data volumes.

Banks are in a unique position to ride this wave. First, banks are one of the most data-rich industries: every transaction, app click, call center call and branch visit leaves a structured digital footprint. Second, banking business margins are highly sensitive to operational efficiency and the quality of risk decisions — two things that are precisely AI's strengths. Third, customer expectations have been shaped by the tech giants: they demand experiences that are instant, personalized and available 24 hours a day, standards that can only be met with intelligent automation.

1.2 "Building the AI-Powered Bank" Vision Map



The nine use cases of the AI-powered bank at a glance

NVIDIA's "Building the AI-Powered Bank" infographic depicts a future bank branch where nine AI capabilities operate simultaneously in one ecosystem. These nine capabilities — which form the framework of Part II of this book — can be grouped into four broad themes:

Customer experience theme. Three capabilities focus directly on interactions with customers: **Personalized Branch Banking** provides branch staff with recommendations for the next best action so that service is fast and personalized; **Avatar Assisted Banking** presents digital avatars powered by conversational AI that serve customers independently; and **Virtual**

Financial Advice allows customers to meet financial advisors via virtual reality or augmented reality in a metaverse environment.

The theme of decision intelligence. Three capabilities work in the financial decision space: **Fraud Alert** uses deep learning to improve fraud detection accuracy and send real-time notifications; **Automated Investment Advisor** uses an algorithm-based robo-advisor to select and manage investment portfolios; and **Data Visualization** strengthens risk management and drives predictive analytics for decision makers.

The theme of physical security. Two capabilities employ computer vision: **Intelligent Video Analytics (IVA) for Threat Detection** analyzes individual behavior within the branch to identify suspicious activity or unknown objects; and **IVA for ATM Security and Monitoring** monitors ATMs for suspicious activity for customer safety and fraud prevention.

Operational and design themes (operations and design). One capability works behind the scenes: **Digital Twins** build digital twins of bank branches to train employees, design new layouts, and test operational changes before they are implemented in the real world.

What's interesting about this vision map is not each capability in isolation — many of which banks are already executing individually — but rather the idea that they all share the same foundation: unified data, accelerated computing infrastructure, and an enterprise AI platform. The branch in the picture is not a collection of nine projects, but rather a single organism with nine organs.

1.3 Why Now: Three Convergent Drivers

The adoption of AI in banking is being accelerated by three forces converging at the same time.

Technology driver. GPU-based accelerator computing reduces model training time from weeks to hours; modern deep learning architectures (especially transformers) surpass classical methods in accuracy on text, sound, and imagery; and generative AI opens up a new class of applications — assistants that can actually have conversations, summarize hundreds of pages of credit documents, or write draft responses to customer complaints. At the same time, the software ecosystem is maturing: open source frameworks, pre-trained models, and MLOps platforms are making capabilities once exclusive to tech giants affordable to mid-sized banks.

Economic driver. Pressure on net interest margins, rising compliance costs and competition from fintechs are forcing banks to cut operating costs while increasing commission-based revenues. AI offers both: automation cuts costs per service transaction, while personalization increases cross-sell and customer retention. On the risk side, small improvements in the accuracy of fraud models or credit models translate directly into material loss savings — for large banks, a reduction in false negative fraud of just a few basis points is worth tens to hundreds of billions of rupiah per year.

Drivers of behavior and regulation. Digital generation customers expect instant and personalized service; they switch banks more easily than previous generations. Regulators, on the other hand, are starting to provide a clearer framework: personal data protection provisions, information technology risk management guidelines, and active discussions on AI governance give banks certainty of direction — as well as new obligations that are most efficiently met with the help of AI (e.g. anti-money laundering transaction monitoring).

1.4 Anatomy of an AI-Powered Bank

Before getting into the use cases, it's important to understand that an "AI-powered bank" doesn't mean a bank with lots of models, but rather a bank whose operating architecture is designed to allow data-driven decisions to flow end-to-end. Anatomy generally consists of five layers:

1. **Data source layer** — core banking, payment switching, digital channels, CRM, security logs, CCTV cameras, IoT sensors in branches and ATMs, as well as external data (credit bureaus, market data, social media as permitted).
2. **Data platform layers** — batch and streaming ingest, lakehouse storage, data catalog and governance, feature store that delivers consistent features for training and inference.
3. **AI platform layers** — data scientist experimentation environment, GPU training cluster, model registry, serving inference (real-time, batch, and streaming), and model monitoring tools.
4. **Intelligent application layer** — nine use cases on top plus others: credit scoring, collections, marketing, document operations, and so on, all consuming services from the AI platform layer via APIs.
5. **Cross-layer governance layers** — cybersecurity, model risk management, regulatory compliance, AI ethics, and human oversight spanning across layers below.

The most common mistake AI startup banks make is building from layer four straight (buying intelligent apps per use case) without solid layers two and three. The result: model silos that can't share features, recurring integration costs, and audit nightmares. The book consistently advocates a platform approach: build a foundation once, use it many times.

1.5 Business Value: The Four Quadrant Framework

To prioritize AI investments, the following simple framework of four value quadrants has proven useful:

Quadrant	Sources of Value	Use Case Example	How to Measure
Revenue grows	Cross-sell, retention, acquisition	Next best action, robo-advisor	Conversion uplift, CLV

Quadrant	Sources of Value	Use Case Example	How to Measure
Costs decline	Automation, process efficiency	Avatar banking, IVA	Cost per interaction, FTE
Losses decline	Fraud, credit losses	Fraud alert, credit scoring	Loss rate, recovery
Controlled risk	Compliance, security	AML monitoring, IVA ATM	Audit findings, incidents

Each use case in Part II will be mapped to this quadrant, complete with performance indicators commonly used to build the business case.

1.6 Reality on the Ground: Between Demo and Production

Industry statistics consistently show the gap between experimentation and production: the vast majority of AI proofs of concept never reach production scale. The cause is rarely algorithmic. The model that won in the data scientist's notebook failed in production because the production data was different from the training data, because integration into the core system took months, because no one thought about who was monitoring the model as it drifted, or because the compliance unit only got involved at the end and discovered the underlying problem.

Therefore, throughout this book, “successful” is defined not as model accuracy in the laboratory, but as a system that runs in production, is monitored, audited, and proven to drive business metrics. The chapters of Part III (MLOps, security, regulation, implementation strategies) exist precisely to close that gap.

1.7 Chapter Summary

- Banking is entering a third wave of transformation: from recording transactions to understanding and predicting, based on systems that learn from data.
- The “Building the AI-Powered Bank” vision map summarizes nine capabilities under four themes: customer experience, decision intelligence, physical security, and design-operations.
- Three converging drivers — technology (GPUs, deep learning, generative AI), economics (margins and competition), and behavioral-regulatory — make AI adoption a strategic imperative, not an option.
- AI-powered banks are about an integrated five-layer architecture, not a collection of disparate model projects.

- Business value is measured on four quadrants: revenue up, costs down, losses down, risk under control; True success is measured in production, not in notebooks.

VISUAL TAKEAWAY

The AI Banking Value Flywheel



CHAPTER 2 — Foundations of Technology: From Machine Learning to Generative AI

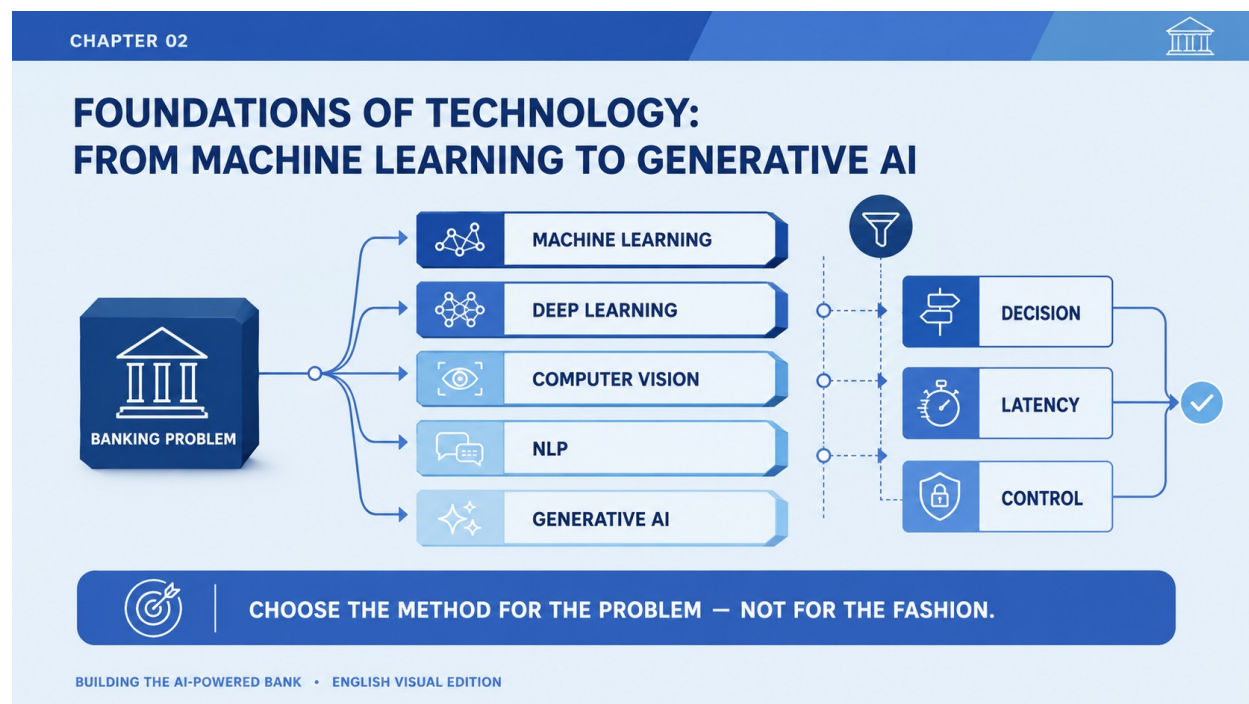


Figure 2. A Practical Map of Artificial Intelligence — Choose the simplest method that satisfies the decision, latency, and control requirements.

2.1 The Big Map of Artificial Intelligence

VISUAL TAKEAWAY Method selection should follow the problem and control environment—not fashion.

The term AI is often used loosely, so let's establish a taxonomy. Artificial Intelligence is a big umbrella: any technique that makes machines exhibit intelligent behavior. In it, machine learning (ML) is the dominant approach today: instead of being programmed with explicit rules, systems learn patterns from data. Within ML, deep learning (DL) is a family of multi-layer artificial neural network-based methods that excel on unstructured data — images, sounds, text. And within DL, generative AI is the newest generation: models that don't just classify or predict, but generate new content — text, images, sounds, code.

For banks, these four taxonomy layers are all relevant and complementary. Classic statistical and ML models (logistic regression, gradient boosting) remain the backbone of credit scoring due to their explainability; deep learning masters modern fraud detection, computer vision, and language processing; Generative AI opens up a class of conversational and productivity applications. Good engineers choose the simplest tools that meet a need, not the most sophisticated ones.

2.2 Machine Learning Paradigms

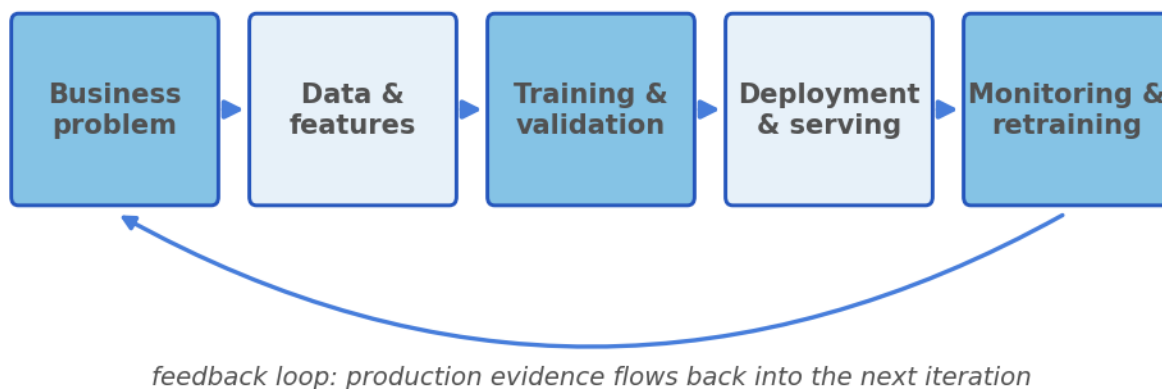
Supervised learning (directed learning) trains a model from input-label pairs. Banks use this the most: transactions labeled fraud/not-fraud, customers labeled churn/remaining, credit applications labeled current/bad. Two families of tasks: classification (categorical labels) and regression (numerical labels, for example estimating customer income).

Unsupervised learning (undirected learning) finds structure in unlabeled data: clustering for customer segmentation, anomaly detection to find "strange" transactions without previous examples of fraud, dimensionality reduction for visualization. This is crucial in banking because fraud and money laundering labels are always rare and late.

Semi-supervised and self-supervised learning make use of abundant unlabeled data. Self-supervised — where a model creates “labels” from its own data (e.g. guessing a hidden word in a sentence) — is the secret behind large language models and modern foundation models.

Reinforcement learning (RL) trains agents through trial and error with reward signals. In banking, RL is explored for optimizing billing strategies, executing orders in trading, and determining the order of product offers; it demands extra caution because agents “learn by acting” in the real world.

2.3 Anatomy of a Machine Learning Project



The machine learning lifecycle: a loop, not a line

Regardless of the algorithm, almost all ML projects follow the same cycle, which every stakeholder needs to understand:

1. **Problem formulation** — translating business objectives into a defined prediction task: what the prediction unit is, what its label is, when predictions are needed, and what decisions will be made from them.

2. Data collection and labeling — pulling historical data, carefully defining labels (when is a legitimate transaction considered fraud?), and dealing with label leakage — classic mistakes when features secretly contain future information.
3. Feature engineering — turning raw data into predictive signals: time window aggregation (number of transactions in the last 1 hour), ratios, entity embeddings, and so on. In deep learning, some of these techniques are learned automatically by the network.
4. Training and validation — splitting the data with discipline (on time-dimensional data, always split by time, not random), selecting algorithms, tuning hyperparameters, and measuring performance on data the model has never seen.
5. Business evaluation — translating statistical metrics into business metrics: how many dollars in losses were prevented, how many good customers were rejected incorrectly, how much analyst workload emerged from alerts.
6. Deployment and monitoring — serving models as a service, monitoring prediction quality and data drift, and setting up retraining mechanisms. This stage is thoroughly dissected in Chapter 15.

2.4 Classic Algorithms that Remain the Backbone

Logistic regression models class probabilities via a sigmoid function over linear combinations of features. It's fast, stable, and each coefficient can be read as a feature contribution — why traditional credit scorecards (based on weight of evidence and score points) are essentially neatly packaged logistic regressions, and why regulators love them.

Decision trees and their ensembles dominate tabular banking data. Random forest averages many trees trained on random samples; gradient boosting (with popular implementations such as XGBoost and LightGBM) builds trees sequentially, each tree correcting the merge errors of the previous trees. On structured transaction and behavioral data, gradient boosting often beats deep learning with much less computational cost — a fact worth remembering before being tempted to use neural networks for everything.

Distance and kernel-based methods (k-nearest neighbors, support vector machine) as well as probabilistic models (naive Bayes, Gaussian mixture models) remain useful in special cases: similarity search, small datasets, or the need for uncertainty estimation.

Classic anomaly detection methods — isolation forest, local outlier factor, one-class SVM — become the first layer of defense when crime labels are not yet available.

2.5 Deep Learning: Building Blocks

Artificial neural networks are composed of neurons: units that compute a weighted sum over their input and then pass it through a nonlinear activation function (ReLU, GELU, sigmoid). Stacked layers of neurons form a network; training is the process of finding weights that minimize the loss function

— for example cross-entropy for classification — through gradient descent: the loss gradient for each weight is calculated by backpropagation, then the weights are shifted in the opposite direction of the gradient. Optimizer variants such as Adam adjust the learning steps adaptively. Regularization techniques — dropout, weight decay, early stopping, batch normalization — keep models from overfitting, that is, memorizing training data instead of learning general patterns.

Three architectural classics form the foundation:

- **Convolutional Neural Network (CNN)** uses convolution filters that sweep the image to detect local patterns (edges, textures, shapes) in a hierarchical manner. CNN is the engine behind banking computer vision: document reading, facial verification, and video analytics in Chapters 9 and 13.
- **Recurrent Neural Network (RNN)** and its successor LSTM/GRU process sequences with internal memory — suitable for transaction series and time signals, although they are now largely replaced by transformers.
- **Autoencoders** learn to compress input into a compact representation and then reconstruct it; Large reconstruction errors signal anomalies — an important technique for fraud and monitoring.

2.6 Transformer: The Architecture That Changes Everything

Introduced in 2017, the transformer architecture replaces recurrence with a self-attention mechanism: each element in a sequence (word, transaction token, image patch) calculates how relevant other elements are to itself, then builds a contextual representation of those relationships. Attention is calculated from three projections — query, key, value — and is carried out in parallel in many "heads" (multi-head attention), interspersed with feed-forward layers, residual connections, and normalization layers.

Two properties make transformers revolutionary. First, parallelism: because there is no recurrence, the entire sequence is processed synchronously — perfect for GPUs. Second, scalability: performance continues to improve as data, parameters and computation (scaling laws) increase, giving rise to giant foundation models. Encoder variants (such as the BERT family) excel for text understanding and classification; decoder variants (GPT family) excel for text generation; vision transformer brings the same paradigm to imagery; and tabular/sequence transformers are starting to be used for customer transaction series.

2.7 Natural Language Processing for Banking

Natural Language Processing (NLP) touches banks at many points: complaint classification, entity extraction from credit documents, sentiment analysis, filtering negative news for due diligence, document summarization, and of course chatbots. The technical chain starts from tokenization (breaking text into subword tokens), embedding (mapping tokens to continuous vectors that capture

meaning), then contextual models (transformers) that generate representations for downstream tasks.

For Indonesian, the practical challenges are real: a mix of languages (Indonesian-English-regional), typical chat abbreviations, and local banking terms. General strategy: start from a pre-trained multilingual model, perform fine-tuning on the domain corpus (customer chat logs, product documents), and build a normalization dictionary for abbreviations. Evaluation must use real conversation data, not standard sentences.

2.8 Computer Vision: Machines that See

Core computer vision tasks relevant to banks: image classification (document types), object detection (finding people, bags, vehicles in frames and their bounding boxes), segmentation (separating object pixels from the background), pose estimation (body joint points for motion analysis), multi-object tracking (following object identities between frames), action recognition (classifying activity in video clips), and OCR (reading text in documents). Modern single-stage detection models achieve a speed-accuracy balance that enables real-time analysis of dozens of video channels per GPU — the technical foundation of Chapters 9 and 13.

An important principle: the performance of a vision model is highly dependent on the suitability of the training data to field conditions — camera angles, lighting, resolution, occlusion. A winning model on a public dataset can fail in a dimly lit ATM aisle; therefore field data collection and augmentation (brightness variations, blur, angles) are an integral part of a vision project.

2.9 Speech AI: Giving Machines Ears and a Voice

Oral conversation involves two-way technology. Automatic Speech Recognition (ASR) converts sound waves into text; Modern end-to-end models (based on conformer/transformer architectures) are trained on thousands of hours of audio and can be fine-tuned for local accents and terms. Text-to-Speech (TTS) does the opposite: generates natural speech from text, now with almost indistinguishable human quality, complete with intonation control. Between the two there are supporting components: voice activity detection, diarization (separating speakers), and noise suppression. These three are the vital organs of avatar banking (Chapter 8) and call center analytics.

2.10 Generative AI and Large Language Models

The Large Language Model (LLM) is a giant decoder transformer that is trained to predict the next token in a massive text corpus, then aligned (aligned) through instruction tuning and reinforcement learning from human feedback to safely follow instructions. Emerging capabilities — understanding long contexts, incremental reasoning, writing in multiple styles,

calling functions/tools — make it a versatile engine for conversational and productivity applications.

For banking, three main utilization patterns:

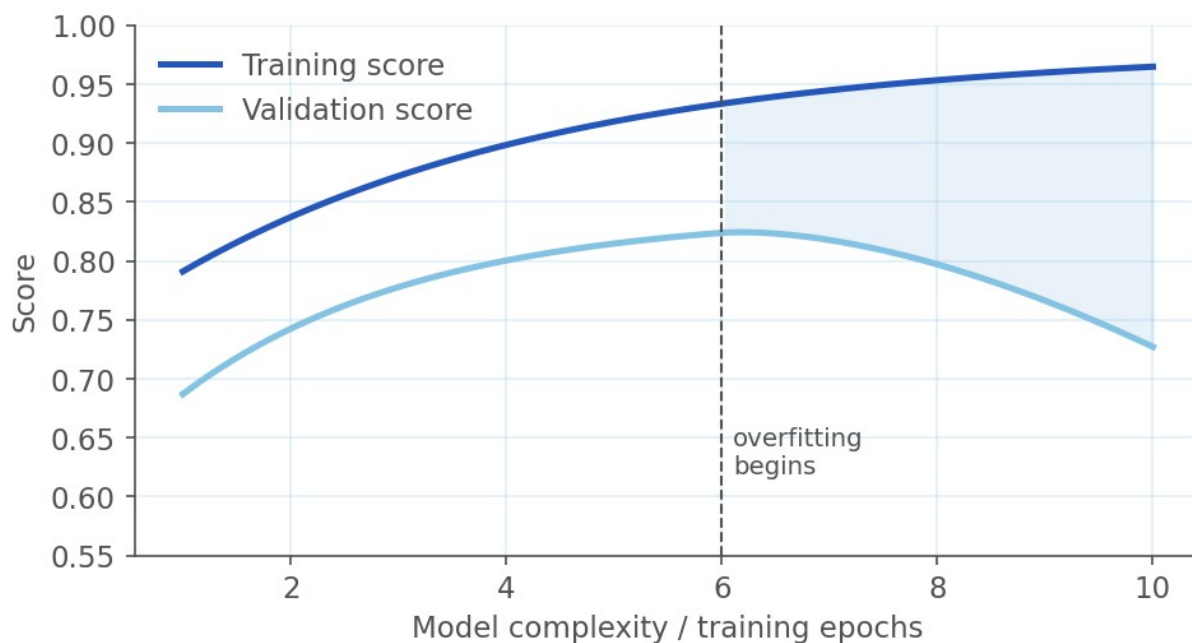
1. **Prompt engineering** — using ready-to-use models with structured instructions; fast, suitable for general tasks.
2. **Retrieval-Augmented Generation (RAG)** — the most important pattern for enterprises: internal documents (policies, products, procedures) are cut into parts, converted into embeddings, stored in vector databases; when a customer asks, the system takes the most relevant snippets and injects them into the prompt, so that model answers are based on authoritative sources and can be cited. RAG suppresses hallucinations and allows updating knowledge without retraining the model.
3. **Fine-tuning** — advanced training of models (often with parameter-saving techniques such as LoRA) on domain data for specific styles, formats, or knowledge; used when prompt and RAG are not enough.

Two inherent risks LLMs must be managed from design: hallucinations (cogent but erroneous answer) — mitigated by RAG, coverage limitations, and citation obligations; as well as prompt injection (malicious input that hijacks instructions) — mitigated by instruction-data separation, input-output filtering (guardrails), and limiting the access rights of model-callable tools. Complete governance is discussed in Chapter 16.

2.11 Graph Neural Networks: Learning from Relationships

Financial crimes rarely stand alone; it lives in a network — accounts that transfer to each other, devices that multiple accounts share, addresses shared by multiple credit applications. Graph Neural Network (GNN) models data as a graph (node = entity, edge = relationship) and learns the representation of each node by aggregating its neighbors' information in layers (message passing). The result: a model that can detect fraud syndicates, multi-layered money laundering networks, and merchant collusion invisible to per-transaction models. GNN will appear as the main weapon in Chapter 11.

2.12 Measuring Performance: Metrics Everyone Should Understand



Overfitting: training and validation scores diverge as complexity grows

The biggest communication mistakes between data teams and the business are usually rooted in metrics. Summary that must be mastered:

Metric	Practical Meaning	When to Use
Precision	Of those flagged positive, how many are correct	The cost of false alarms is high
Recall (TPR)	Of all positive cases, how many were detected	The cost of missed positives is high
F1-score	Harmonic mean of precision and recall	A balance is required
ROC-AUC	Ability to rank positives above negatives	General model comparison
PR-AUC	Similar, but more informative under class imbalance	Fraud, AML (rare positives)

Metric	Practical Meaning	When to Use
KS / Gini	Risk-score discrimination	Credit scoring
Calibration	A score of 0.8 should mean 80%	Scores used as probabilities

Two additional disciplines: on extreme lame class problems (fraud is often below 0.1% of transactions), total accuracy is almost meaningless — models that answer “not fraud” for everything are 99.9% accurate but useless; and each metric must be evaluated at critical slices — per customer segment, per channel, per region — to catch any inequities or blind spots hidden in the average.

2.13 From Model to System: Latency, Throughput, and Cost

Models are just components; the bank bought the system. Three engineering quantities determine production feasibility. Latency — the time of one prediction — is limited by business requirements: card transaction authorization typically requires fraud decisions in tens of milliseconds; chatbots tolerate hundreds of milliseconds to several seconds. Throughput — predictions per second — determines the number of instances and GPUs needed at peak load (think payday night or a flash sale). Cost per million predictions combine the two with infrastructure prices. Inference optimization techniques — quantization (reducing the numerical precision of weights, e.g. FP16/INT8), pruning, distillation of large models into small models, dynamic batching, and graph compilation — can reduce costs by many orders of magnitude without sacrificing accuracy significantly, and will be discussed operationally in Chapters 3 and 15.

2.14 Chapter Summary

- Taxonomy: AI $\supset$ machine learning $\supset$ deep learning $\supset$ generative AI; choose the simplest tool that meets your needs.
- The four learning paradigms (supervised, unsupervised, self-supervised, reinforcement) are all used in banking; Rare and late labels make no-label methods very valuable.
- Gradient boosting masters tabular data; CNN owns the image; transformers master language, sound, and increasingly penetrate transaction sequences; GNN masters relationships between entities; autoencoder masters anomalies.
- LLM is leveraged through prompts, RAG, and fine-tuning; hallucinations and prompt injection must be mitigated from design.
- Master metrics (precision, recall, PR-AUC, KS, calibration) and engineering quantities (latency, throughput, cost) — the shared language between data scientists, engineers, and business.

2.15 Technical Insert: Mathematical Intuition in Five Concepts

Readers don't need to be mathematicians to navigate an AI project, but the following five intuitions prevent the most costly misunderstandings.

Probability as an output language. A good model score is a calibrated probability: of all cases scoring 0.8, about 80% are truly positive. Calibration allows scores to be used in decision calculations (expected value = probability $\times$ impact) and compared across models; without calibration, "0.8" is just a rating number. Post-training calibration techniques (Platt scaling, isotonic regression, temperature scaling) are cheap and almost always feasible.

The loss function is a contract. The model optimizes exactly what it says in its loss function — nothing more. If the business cost of a false negative is ten times a false positive, the contract must go to loss weighting or to threshold determination; expecting the model to “figure it out on its own” is a classic source of disappointment.

Bias-variance and data size. Simple models can underfit (high bias: pattern not captured); complex models can overfit (high variance: memorizing noise). The cure for overfitting is not just regularization, but primarily more and more diverse data — and the learning curve (performance vs amount of data) is the cheapest diagnostic tool to decide whether the next investment should be in the data or in the architecture.

Correlation, causality, and intervention. Predictive models learn correlation; Business decisions often require causality ("if we make an offer, will customers buy because of that offer?"). Causal answers demand experimental designs (randomized control groups) or causal methods — this is at the root of why uplift modeling (Chapter 5) differs fundamentally from ordinary propensity, and why the book repeatedly calls for control groups.

The high dimensions are strange. Over hundreds of features, everyday geometric intuition breaks down: points become equally “distant”, and distance-based models weaken. Deep learning overcomes this by learning meaningful lower-dimensional representations (embeddings) — the reason why embeddings are the universal currency of modern AI: customers, products, documents, and images all end up as vectors that can be compared.

VISUAL TAKEAWAY

Five Concepts, One Production Discipline

1

Probability

calibrated confidence

2

Optimization

minimize loss

3

Generalization

work beyond training

4

Correlation

not causality

5

Bias-variance

fit without fragility

CHAPTER 3 — Accelerated Computing Infrastructure: GPUs, Cloud, and Edge

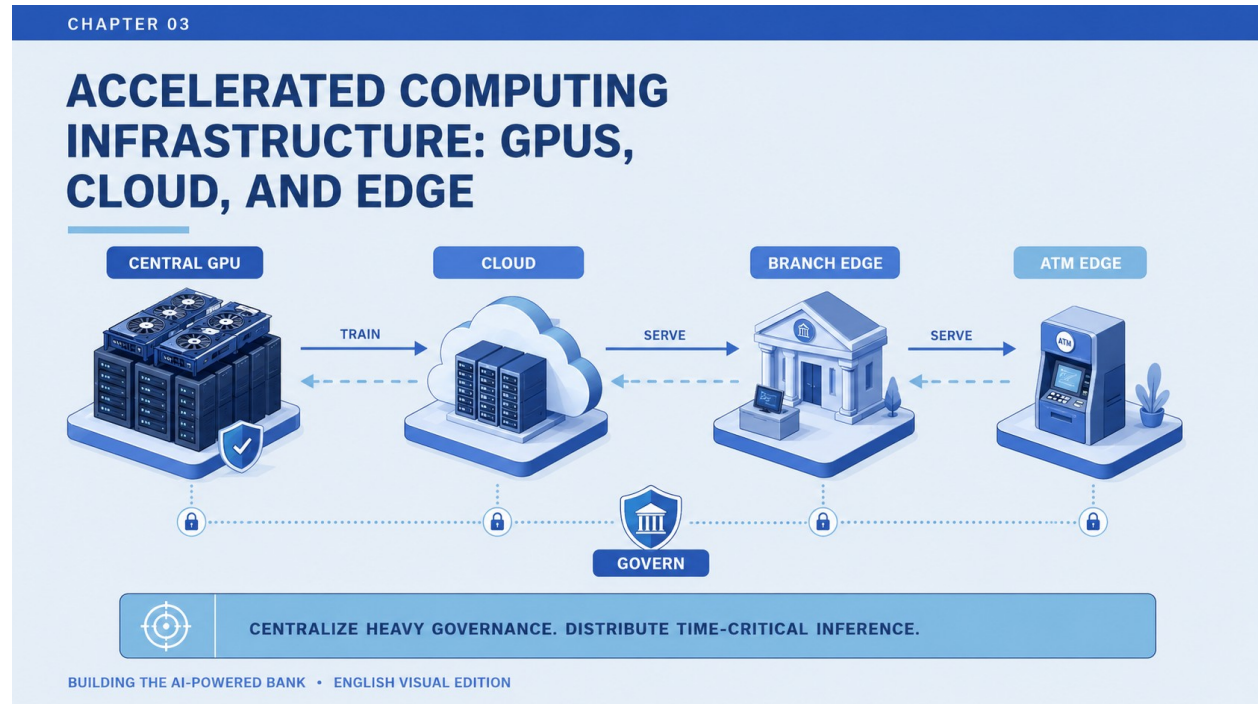


Figure 3. Hybrid AI Deployment Topology — Place workloads where latency, data gravity, resilience, and regulation make the most sense.

3.1 Why GPUs Are the Heart of AI

VISUAL TAKEAWAY A hybrid topology keeps heavy governance centralized while moving time-critical inference closer to the event.

CPUs are designed to execute a small stream of instructions no matter how complex they are very quickly; GPUs are designed to execute thousands of simple operations synchronously. A world-changing coincidence: deep learning's core workloads — matrix multiplication and convolution — take the form of exactly thousands of identical simple operations. A single modern data center GPU contains thousands of parallel cores plus dedicated Tensor Core units that execute mixed-precision matrix multiplication (FP16/BF16/FP8) in a single cycle, supported by high-bandwidth memory (HBM) that feeds data to the cores at rates of hundreds of gigabytes to terabytes per second. The practical result: training that would take weeks on a CPU cluster is completed in a matter of hours on multiple GPUs, and real-time impossible inference becomes routine.

Its software ecosystem is as important as its silicon. CUDA is a parallel programming platform that is the foundation of almost all deep learning frameworks; on top of it stand the optimized libraries (cuDNN for neural network operations, cuBLAS for linear algebra, NCCL for inter-GPU

communication) that make PyTorch and TensorFlow automatically run fast. For bank engineers, the implication is simple but important: the choice of AI infrastructure is not just a matter of hardware specifications, but rather the depth and maturity of the software stack on top of it.

3.2 Data Center GPU Platform Anatomy

For serious workloads, GPUs do not stand alone but are assembled:

- **Multi-GPU servers connect 4-8 GPUs with ultra-high-speed interconnects (NVLink/NVSwitch) so that the GPUs can exchange gradients and activations almost like one shared memory — crucial for training large models that cannot fit on a single GPU.**
- **Multi-node clusters chain together tens to thousands of servers over a low-latency network (InfiniBand or Ethernet RoCE) with technologies like GPUDirect RDMA that move data between GPUs across servers without bypassing the CPU.**
- **High-throughput parallel storage (NVMe, parallel filesystem) keeps the GPU from being "data starved"; in vision training and LLM, the bottleneck is often in the data pipeline, not computing.**
- **Schedulers and orchestrators — Kubernetes with operator GPUs, or HPC schedulers (Slurm) — divide clusters across teams with quotas, queues, and isolation.**

The GPU partitioning feature is worth recognizing: Multi-Instance GPU (MIG) partitions a single physical GPU into multiple isolated instances with guaranteed memory and bandwidth — ideal for inference of many small models or multi-tenant environments; while time-slicing shares GPUs in turns for light loads such as experimental notebooks.

3.3 Training Large-Scale Models: Parallelism Strategy

When the model or data extends beyond a single GPU, training is split:

1. **Data parallelism — each GPU holds a copy of the full model and processes different batch fractions; Gradients are averaged across GPUs (all-reduce) each step. Standard strategy for the majority of banking models.**
2. **Model/tensor parallelism — layers or weight matrices split across GPUs; used when a single model does not fit in the memory of one GPU.**
3. **Pipeline parallelism — different groups of layers are placed on different GPUs and batches are flowed in stages like a conveyor belt.**
4. **Memory-efficient techniques — mixed precision (FP16/BF16 with loss scaling), gradient checkpointing (exchanging compute for memory), and sharding state optimizers (ZeRO/FSDP family) — allow much larger models to be trained on the same device.**

For banks, the need to train LLMs from scratch is almost non-existent; a realistic one is fine-tuning the foundation model (multiple GPUs, hours-to-days) and training an internal tabular/vision/graph model (one to a few nodes). Cluster capacity should be planned from this real load, not from fantasy.

3.4 Production Inference: The Art of Serving Predictions

Inference has a different engineering profile than training: tight latency, fluctuating load, and running 24×7. The default architectural pattern is a centralized inference server — the canonical example being the Triton Inference Server — that loads multiple models from a repository, serves requests via HTTP/gRPC, and provides production capabilities: dynamic batching (merging requests that arrive close together for GPU efficiency), concurrent execution of multiple models on a single GPU, inter-model ensembles/pipelining, sequence model-specific scheduling, and Prometheus-ready metrics.

In the model optimization layer, compilation and quantization (e.g. TensorRT for vision/tabular, and LLM-only runtime with paged attention, KV-cache, and in-flight batching techniques for language models) commonly provide multiple speedups: FP8/INT8 precision on Tensor Core, kernel fusion, and memory structuring. Best practice: set a latency budget per use case (e.g. $p99 \leq 50$ ms for fraud scoring, ≤ 300 ms for dialog components), then optimize backwards from there — measure, quantize, batch, then add hardware.

3.5 Edge AI: Intelligence in Branches and ATMs

Not all inferences are worth sending to the data center. Video analytics from 24 cameras per branch would flood the WAN if raw video was sent continuously; ATM security decisions must continue when the connection is lost; and regulations could require the video not to leave the location. The answer is edge computing: AI-powered computing devices installed on-site — from embedded-class modules (Jetson family) inside ATMs to GPU-enabled edge servers in branch server rooms — running models directly on-site and sending only metadata/events to the center.

Typical edge challenges that must be planned for from the start: fleet management of thousands of devices — provisioning, over-the-air firmware and model updates, health monitoring; power and thermal limitations (models must be quantized and trimmed to fit); physical security of the device (secure boot, disk encryption, device-identity binding); as well as drop-off operations (store-and-forward when the network is down). Mature hybrid architectural pattern: edge handles real-time detection, center handles retraining, cross-site aggregation, and long-term analytics.

3.6 Cloud, On-Premise, or Hybrid

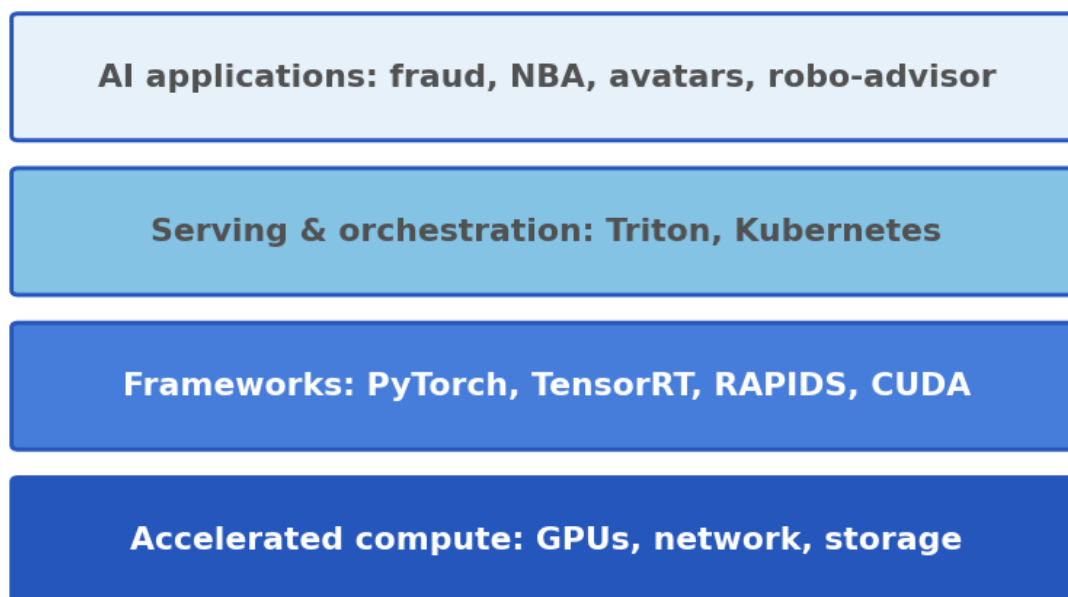
AI workload placement decisions in banks are pulled by four forces: regulation and data residency, economics, speed, and talent availability.

Criteria	Public Cloud	On-Premise	Hybrid
Start time	Days	Months	Phased

Criteria	Public Cloud	On-Premise	Hybrid
Elastic scale	Excellent	Limited	Good
Data control	Requires strict controls	Full	Segmented
Biaya beban tetap 24x7	Often expensive	More efficient	Optimal
Research load fluctuates	Excellent fit	Inefficient	Suitable
Local compliance	Requires assessment	Easier	Case by case

The most successful patterns in banks: data classification-based hybrids — experimentation and training on anonymized/tokenized data can use cloud elasticity; inference that touches sensitive personal data and core systems run on-premises or in the on-premises cloud with bank-owned key controls. Either way, two disciplines are mandatory: portability (containers, IaC, storage abstraction) to avoid lock-in, and FinOps (per-project cost tagging, quotas, autoscaling down, GPU utilization monitoring) — idle GPUs are the most common budget leak in AI programs.

3.7 Enterprise AI Software Platform



Each layer only performs as well as the layer beneath it

The accelerated computing stack, from hardware to applications

On top of the hardware, banks need a consistent platform layer: a container registry with curated and vulnerability-scanned framework images; pre-trained model catalogue; domain toolkits (vision, speech, LLM) that accelerate development; and enterprise support (security patches, long-term compatibility) — needs that are addressed by both commercial distributions such as NVIDIA AI Enterprise and open source builds maintained by internal platform teams. His selection principles for banks: software supply chain security (SBOM, image signing), long-term version support and licensing clarity are more important than fringe features.

3.8 Designing Capacity: A Practical Framework

Auditable capacity calculation steps:

1. Workload inventory: model list (type, size), call patterns (average and peak per second), latency budget, retraining schedule.
2. Unit benchmarks: measure throughput per GPU for each model in target optimization configuration (batch, precision) — do not use brochure numbers.
3. Calculate peak requirements with headroom (30–50% typical) and N+1 redundancy per zone.

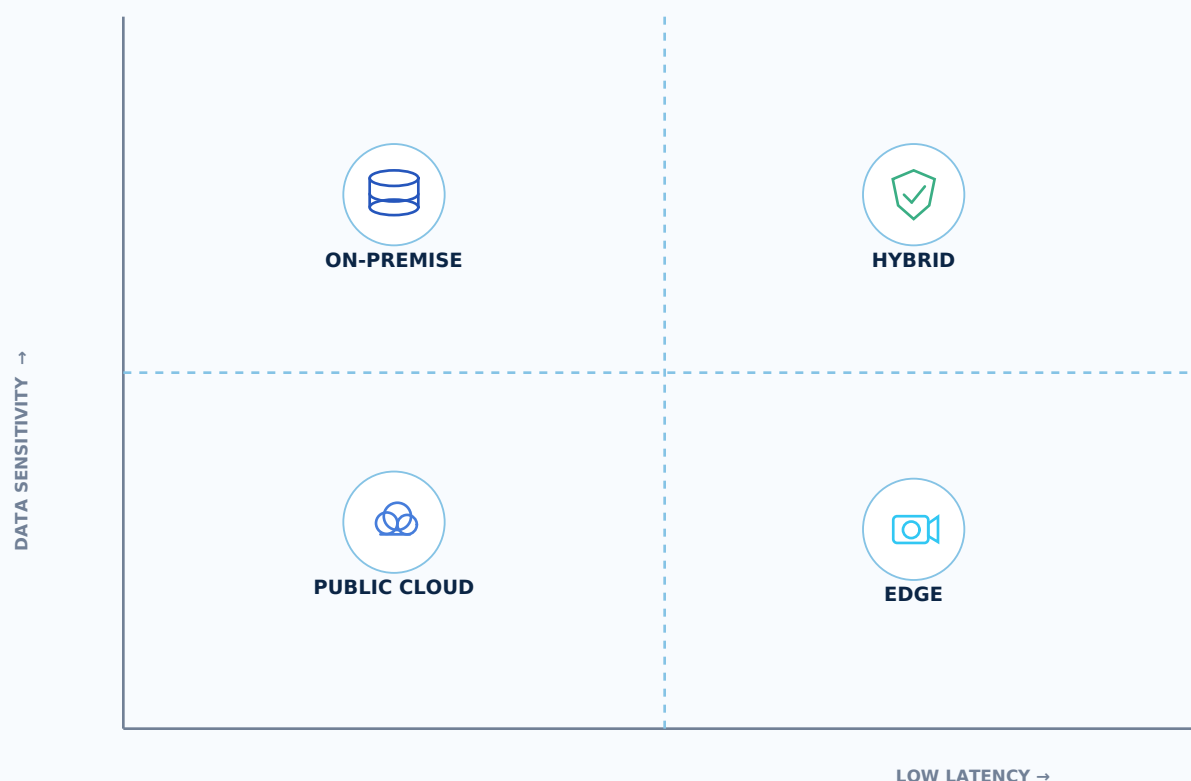
4. Separate pools: training/experimentation (can queue) vs inference (cannot queue), with quotas per team.
5. Plan for 18–24 months of growth and replenishment decision points; review monthly utilization.

3.9 Chapter Summary

- The GPU + CUDA ecosystem changes the economics of AI: training from weeks to hours, real-time inference to routine; Tensor Cores and mixed precision are the key.
- Scalability is achieved via in-server NVLink/NVSwitch circuits, inter-server RDMA networking, and parallelism strategies (data, tensor, pipelines) plus memory-saving techniques.
- Production inference demands an inference server with dynamic batching and model optimization (quantization, compilation); set a latency budget and then backward optimize.
- Edge AI brings intelligence to branches and ATMs; Its success is determined by fleet management, not just models.
- Select cloud/on-prem/hybrid deployment based on data classification and workload economics; maintain FinOps portability and discipline.

VISUAL TAKEAWAY

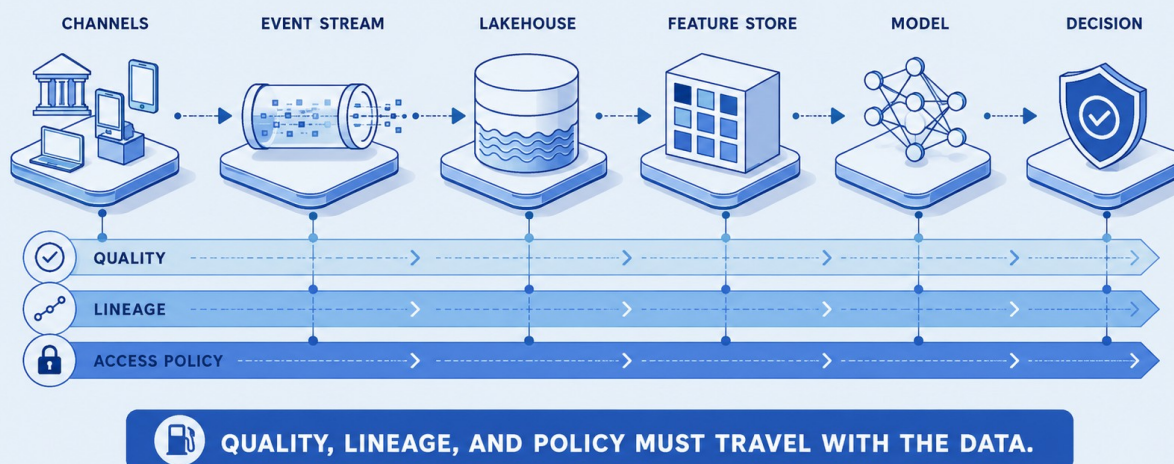
Place Workloads by Latency and Data Sensitivity



CHAPTER 4 — Data: The Fuel of Smart Banking

CHAPTER 04

DATA: THE FUEL OF SMART BANKING



BUILDING THE AI-POWERED BANK • ENGLISH VISUAL EDITION

Figure 4. From Banking Data to Governed Decisions — A modern data foundation connects events, reusable features, models, and accountable action.

4.1 Banking Data Landscape

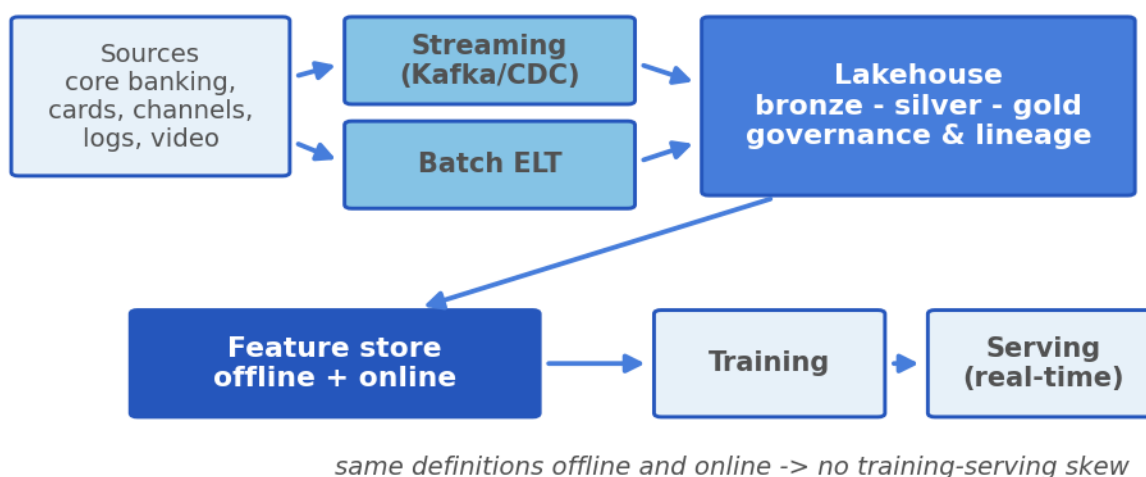
VISUAL TAKEAWAY Quality, lineage, and access policy must travel with the data from source to decision.

The richness of bank data is also its complexity. Key sources to tame:

- Core banking and product systems — accounts, balances, transfers, loans, deposits; source of transaction truth, usually decades old with legacy schemes.
- Switching and payments — card authorization, interbank transfer, QRIS, e-wallet; characterized by very high volume streaming and critical latency.
- Digital channels — mobile and web app clickstreams, session logs, devices, geolocation (as permitted); the raw materials of personalization and behavioral anomaly detection.
- CRM and services — customer profiles, call center interactions (audio and transcripts), complaint tickets, offer history.
- Risk and compliance data — credit bureaus, sanctions and PEP lists, compliance reports, fraud investigation results (expensive gold labels).
- Physical data — branch and ATM CCTV videos, ATM/EDC machine logs, IoT sensors (door, temperature, queue).
- External data — market data, exchange rates, news, demographic data and alternative data as required.

Three prominent properties of this landscape define the architecture: extreme volume and speed at checkout lines; high legal sensitivity (almost all personal data); and historical fragmentation — the result of mergers, legacy systems, and organizational silos.

4.2 Modern Data Architecture: Lakehouse



Reference data flow: sources, lakehouse, feature store, training and serving

Two legacy paradigms complement each other's weaknesses: data warehouses excel at structured analytical queries and governance, but are clunky and expensive for raw, unstructured data; Data lakes are cheap and flexible to accommodate anything, but are prone to becoming “data swamps” with no quality guarantees. The lakehouse architecture combines both: low-cost object storage + open table formats (Delta Lake, Apache Iceberg, Hudi) that add ACID transactions, schema evolution, time travel, and quality enforcement on top of Parquet files.

A proven organizing pattern is the medallion architecture:

1. **Bronze** — raw data as is from the source, unaltered, for auditability and replay.
2. **Silver** — data that is cleaned, deduplicated, normalized, and identity unified (cross-system customer entity resolution — hard work that is worth a lot of money).
3. **Gold** — ready-to-consume tables per domain: customer 360 features, transaction aggregate, risk mart; source for BI and feature stores.

For large organizations, data mesh principles complement the lakehouse: data ownership per business domain (payments, credit, channels) with “data as a product” — each domain is responsible

for the quality, documentation, and SLAs of its data — on top of a self-service platform and centralized federated governance.

4.3 Streaming Pipeline: The Real-Time Nervous System

Use cases such as fraud alerts and next best action depend on data freshness. The backbone is a distributed logging platform (Apache Kafka and similar): every event — card authorization, login, transfer — is published to a topic and can be consumed by multiple systems independently with per-partition order guarantees. On top of it, stream processing (Flink, Spark Structured Streaming, ksqlDB) computes time window aggregation (number of last 5 minutes of transactions per card), stream merging, and pattern detection, with exact-once semantics important for feature consistency.

A key architectural pattern for real-time AI: event-driven feature computation — streaming features are calculated once in the pipeline and written to an online feature store, rather than recomputed by each application. This ensures the fraud model and recommendation system are seeing the same numbers, and makes the feature logic auditable in one place.

4.4 Feature Store: The Bridge Between Data and Models

The feature store is a layer that defines, computes, stores, and serves features consistently for two worlds: offline (historical tables for training, with point-in-time join capabilities that prevent future leaks) and online (millisecond latency key-values for inference). The concrete benefits: eliminate training-serving skew (features in production are different than in training — the most common source of ML bugs), encourage cross-team feature reuse, and provide auditors with a documented feature catalog.

Implementation discipline: every feature has an owner, declarative definition, quality testing, and retention policy; derivative features of personal data recorded for their purpose (purpose limitation); and online-offline materialization originates from the same logic (one definition, two presentations).

4.5 Data Quality and Observability

The best model is defeated by the worst data. The interoperable quality framework includes six dimensions — completeness, uniqueness, validity, consistency, accuracy, timeliness — that translate into automated data tests in the pipeline (a Great Expectations/dbt tests-like framework): value constraints, distribution, referentiality, freshness. On top of it, data observability monitors pipeline metric anomalies (volume suddenly drops 40%? schema changes? distribution shifts?) and links them to derivatives via data lineage — a column-by-column lineage map that answers the auditor's question: "where did this score come from?" Data quality incidents should be treated like production incidents: there's detection, there's escalation, there's post-mortem.

4.6 Governance and Protection of Personal Data

Almost all valuable bank data is personal data, so governance is not bureaucratic but a prerequisite for an operating license. Operational pillars:

- **Classification and catalogue** — each dataset is assigned a sensitivity label (public, internal, confidential, private data, specific personal data) and listed in the catalog with the owner and processing purposes.
- **Basis of processing and limitation of purposes** — in line with the Personal Data Protection Law: data is collected for a specific purpose and is not used beyond it; consents are managed granularly (e.g. personalization permissions are separate from operational permissions).
- **Attribute-based access control** — access is determined by a combination of role, purpose, and sensitivity; dynamic dimasking sensitive columns; every access is recorded for audit.
- **Minimization techniques** — pseudonymization and tokenization (replacing card PAN and NIK with tokens from the point of entry), encryption at rest and in motion with centralized key management, and anonymization/aggregation for experimental environments.
- **Data subject rights and retention** — mechanisms for fulfilling access/deletion requests, retention schedules per category, and self-evident deletion (including in copies and backups).

For sensitive cross-border collaboration and training, privacy-enhancing techniques are worth mastering: federated learning (the model is trained on data locations, only weight updates are exchanged), differential privacy (calibrated noise that limits the leakage of individual information), synthetic data (similarly distributed artificial data for development and testing), and confidential computing (TEE) for processing encrypted data in memory. Each has a utility-privacy trade-off that must be tested, not assumed.

4.7 Unstructured Data: Video, Sound, Documents

The book's nine use cases drag banks into unstructured data territory where governance often lags behind:

- **CCTV video** — huge volume (one branch can generate hundreds of GB/day); practical policy: retain short raws at the edge, keep long only event clips + structured metadata (not faces/identities unless there is a legal basis), and blur/redact for secondary uses like model training.
- **Call center audio** — automatic transcription opens up analytics, but recordings contain specific data (health, financial); apply automatic redaction of sensitive entities to transcripts.
- **Documents** — ID cards, pay slips, financial statements, contracts; AI document pipeline (OCR + classification + extraction) turns it into structured data with confidence scores and a human verification pipeline.
- **Embedding and vector database** — vector representation of text/image for semantic search and RAG; treat embedding personal data as personal data (it may leak its provenance information).

4.8 Integrated Data-AI Reference Architecture

Putting this whole chapter together, the data flow:

1. Source (core, switching, channel, CCTV, call center) →
2. Batch ingest + streaming (CDC from core database; Kafka for events) →
3. Lakehouse medallion (bronze→silver→gold) with catalogue, lineage and quality tests →
4. Feature store (offline for training; online for inference) + vector store for RAG →
5. Consumers: model training, real-time inference, BI/visualization, intelligent applications →
6. Feedback: results of decisions and investigation labels return to the lakehouse (learning loop).

Around that flow stretches cross-cutting controls: IAM and ABAC, encryption and tokenization, quality monitoring, and a comprehensive audit trail. This architecture will be the "ground floor" for each use case diagram in Part II.

4.9 Chapter Summary

- Bank data is rich but fragmented and highly sensitive; the architecture must serve both batch and real-time simultaneously under tight governance.
- Lakehouse with a medallion pattern (bronze-silver-gold) plus data mesh principles balances flexibility, quality, and domain ownership.
- Kafka and stream processing form a real-time neural system; Streaming features are calculated once and presented via the feature store to be consistent and auditable.
- The feature store eliminates training-serving skew — the most common source of ML bugs — and becomes a feature catalog for auditors.
- Personal data governance (classification, purpose limitation, tokenization, retention) plus privacy-enhancing techniques (federated learning, differential privacy, synthetic data) are prerequisites, not complements.
- Unstructured data (video, audio, documents) demands specific policies: event-based minimization, redaction, and retention.

PART II — NINE AI-POWERED BANKING USE CASES

This section dissects one by one the nine capabilities in the "Building the AI-Powered Bank" vision map. Each chapter uses a uniform framework — business context, technical workings, reference architecture, data and models, metrics, implementation stages, risks and mitigations, ending with a checklist — so that it can be directly used as a feasibility study template in your organization.

PART II VISUAL MAP

Nine Use Cases - One Shared Foundation



Personalized branch



Virtual advice



Digital twins



Avatar banking



Video analytics



Risk visualization



Fraud detection



Investment advisor



ATM security

Shared: data • identity • integration • MLOps • controls

CHAPTER 5 — Personalized Branch Banking: Fast Service with Next Best Action

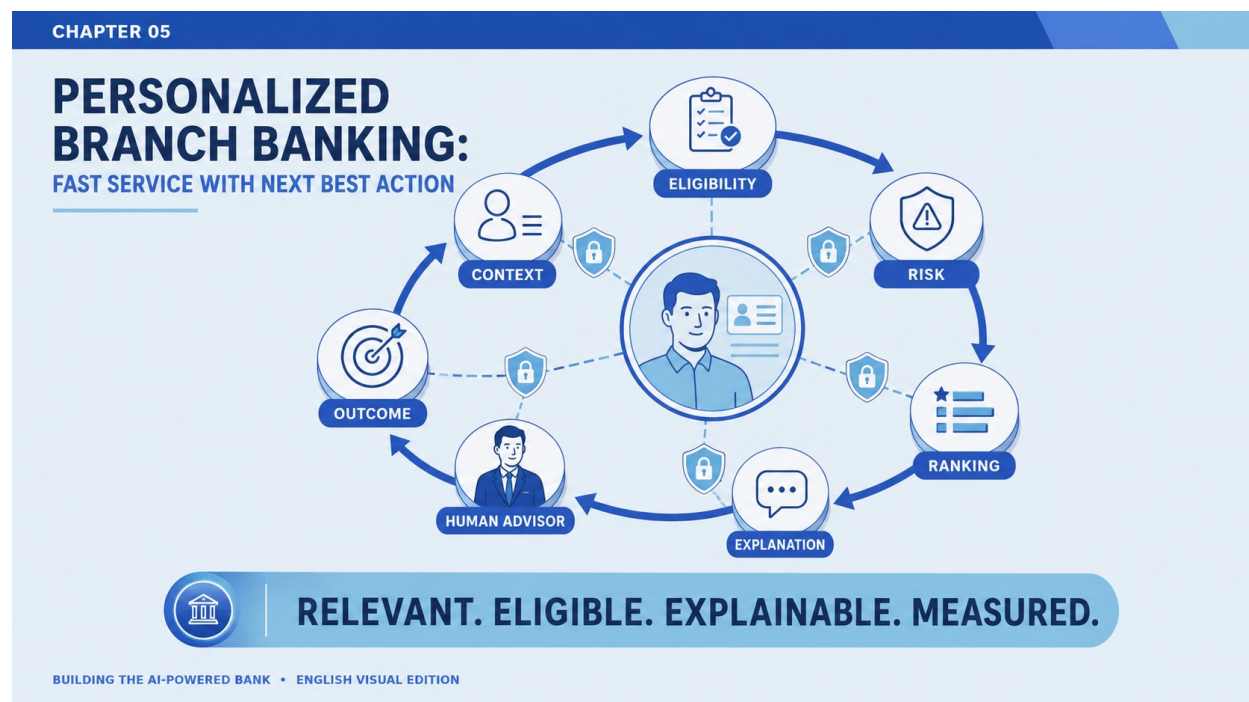
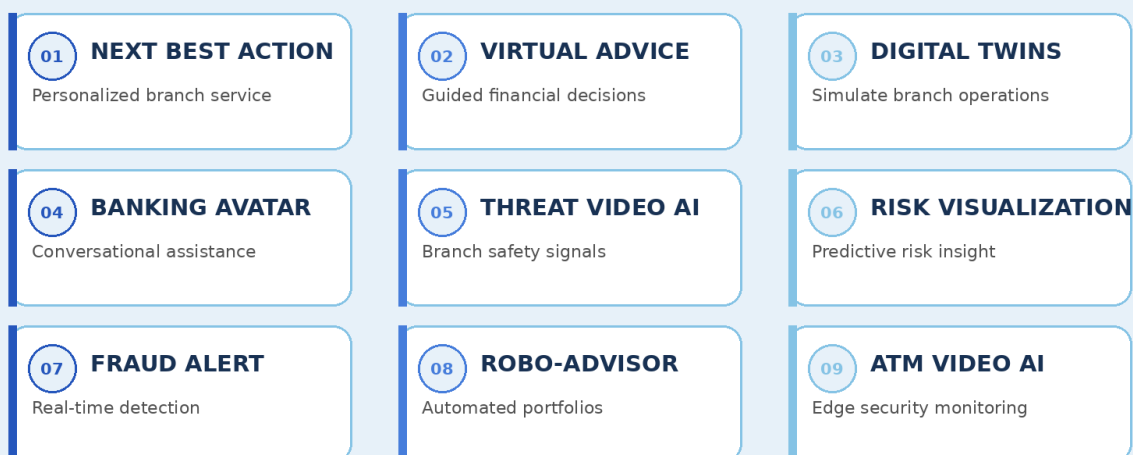


Figure 6. Next Best Action Decision Loop — Context becomes a recommendation only after eligibility, risk, and customer-intent controls.

FIGURE 05

Nine AI-Powered Banking Use Cases

A balanced portfolio across experience, decision intelligence, operations, and physical security.



BUILDING THE AI-POWERED BANK • ENGLISH VISUAL EDITION

Figure 5. Nine AI-Powered Banking Use Cases — A balanced portfolio across experience, decision intelligence, operations, and physical security.

VISUAL TAKEAWAY Portfolio value comes from shared capabilities: data, identity, model operations, observability, and governance.

5.1 Business Context

Imagine two experiences at the branch. First: customers queue, re-explain their identity and needs from scratch, then are offered irrelevant generic products. Second: once a customer checks in, the officer has seen a summary of the customer relationship, the context of the last visit, and the three most relevant recommendations for action — resolving a blocked card, offering partial repayment of the mortgage in favor of the customer, or simply offering nothing because the customer has just filed a complaint. Personalized Branch Banking capabilities turn the first experience into the second: fast, personalized service, driven by next best action (NBA) recommendations — the next best action that a machine calculates for each customer at every moment.

The business value is in two quadrants at once. In the revenue quadrant: relevant offers increase cross-sell/up-sell conversion rates many times over mass campaigns, and the right service moments increase retention. In the cost quadrant: service time per customer is shortened because officers do not fumble, and marketing becomes precise (the promotional budget is directed to customers with the highest propensity). What's often forgotten: NBA is also a customer protection tool — the same system can recommend "do not offer credit" to customers who indicate financial difficulties.

5.2 Technical Workflow

VISUAL TAKEAWAY A useful recommendation is relevant, eligible, explainable, and measured against customer outcomes—not only conversion.

The modern NBA is an orchestration of four components of intelligence:

First, 360 customer profile and segmentation. The foundation is a unified identity (one customer = one entity across products and channels, the result of entity resolution in lakehouse) which is summarized into hundreds of features: demographics, product ownership, transaction behavior, channel interactions, life signals (changes in salary patterns, transactions related to marriage/birth/moving house), and service status (open complaints, arrears). Clustering-based segmentation provides a big map, but the true NBA works at the individual level.

Second, the propensity model. For each candidate action (offer product Mature practice moves from pure propensity to uplift modeling: estimating the difference in outcomes between "given action" and "given no action", so that the budget is directed to customers who are truly affected (persuadables), not those who would buy on their own. Uplift training data

should ideally come from a randomized experiment (control group) designed early in the program.

Third, recommendation engine. For large catalogs of actions, recommender system techniques are used: collaborative filtering and matrix factorization study the pattern of "similar customers choosing similar products"; two-tower neural network maps customers and products to the same embedding space for fast candidate retrieval; Sequence models (transformers over interaction history) capture time dynamics — what is relevant after a new customer opens an account is different from after he or she has paid off a loan. The industry standard pipeline runs two stages: retrieval (distilling thousands of candidates down to tens, of milliseconds) and then ranking (a feature-rich model carefully assesses those dozens of candidates).

Fourth, the decisioning and arbitration layer. Raw scores are not decisions. This layer combines the propensity/uplift score with the business value of each action, eligibility and compliance rules (customer has not met product requirements, contact frequency limits, prohibitions on offering credit to vulnerable segments), channel context, and exploration strategy. For continuous learning, contextual bandit is widely used: the system occasionally tries non-best actions in a controlled manner (exploration) to continuously update estimates, balancing learning and exploiting. The final output: a short list of ordered actions, each with a rationale that the officer can read ("recommended because forex transaction patterns have increased in the last 3 months").

5.3 Reference Architecture



outcomes and front-line feedback are logged as labels for the next training cycle

The next best action decision flow

Runtime flow when the customer arrives at the branch:

1. **Context trigger** — customer check-in (digital queue, card, or recognition with approval); events are published to the event bus.
2. **Feature retrieval** — the NBA service pulls online features from the feature store (profiles, real-time aggregates such as last 24 hours of activity) in a few milliseconds.
3. **Scoring** — inference server executes retrieval + ranking + uplift for a catalog of actions; Typical total latency budget $\leq 200\text{--}300$ ms.
4. **Decisioning** — the rules/arbitration engine applies feasibility, compliance, contact limits, and exploration strategies; generate 1–3 recommendations + reasons.
5. **Presentation** — teller/RM application displays recommendation cards; officers accept, delay, or reject with reasons — valuable feedback.
6. **Results loop** — action results (accepted/rejected/converted in N days) flow back to the lakehouse as labels for subsequent training and uplift measurements.

The same architecture serves other channels (mobile, call center, ATM) — this is why NBA should be built as a centralized decision service, not a feature in a single app: one brain, many faces, with cross-channel orchestration (don't offer in the branch what was just rejected in the app yesterday).

5.4 Data and Model

Realistic minimum data requirements: transaction history and product ownership $\geq 12\text{--}24$ months; channel interaction log; campaign history and results (most often lost — start recording now); and marketing approval markers per customer. Typical models: gradient boosting for per-product propensity (robust and explainable baseline), two-tower + sequence transformer for large-scale recommendations, meta-learner (T/X-learner) for uplift, and contextual bandit (e.g. LinUCB/Thompson sampling) in the arbitration layer. All scores are calibrated to be comparable across models.

5.5 Success Metrics

Levels	Metric	Notes
Model	AUC, calibration, recall@K	Per product & per segment
Uplift	Qini/uplift@decile	Need a control group
Adoption	Staff recommendation adoption rate	Trust indicator
Business	Incremental conversion, revenue/contact	Measure against a control, not the past

Levels	Metric	Notes
Experience	Branch NPS/CES, service time	Avoid turning service into a sales machine
Compliance	Contact-limit breaches = 0	Automatically audited

The golden rule of measurement: always keep a small randomized control group that does not receive NBA; without it, uplift claims cannot be justified.

5.6 Implementation Stages

1. Foundation phase (months 1–3) — unify customer identity, build 100–200 core features in feature store, start recording campaign results with control group.
2. Pilot phase (months 3–6) — 3–5 high-value actions, propensity model + eligibility rules, pilot in a handful of branches with officer training; measure adoption and uplift.
3. Expansion phase (months 6–12) — add action catalog, ranking and arbitration layers, cross-channel orchestration, recommendation rationale in UI.
4. Optimization phase (12+) — full uplift modeling, bandit for exploration, personalization of offer content (including generative AI assistance for message variations — with compliance team approval of templates).

5.7 Risks and Mitigation

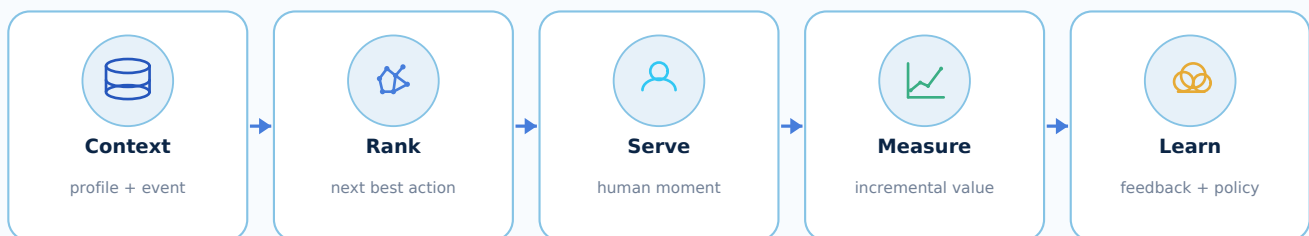
- Offer fatigue and broken trust — limit the frequency of cross-channel contact; include non-sales actions (service, education) in the catalog; give officers a recorded veto.
- Bias and exclusion — audit of recommendations per protected group; ensure the model does not systematically offer expensive products to vulnerable groups; document prohibited features.
- Violation of privacy/consent — enforce purpose restrictions in the feature store: data-driven features without marketing permission should not flow to the marketing NBA.
- Officers ignore the system — involve officers from design, show rationale, and close the loop (show that their feedback changes recommendations).
- Cannibalization and pseudo-metrics — incremental measurement with control; beware of models that are "successful" only because they target people who will definitely buy.

5.8 Chapter 5 Checklist

- Integrated customer identity and online-offline store features are available.
- Every action has a definition of success, feasibility, and business owner.

- Permanently run randomized control group; uplift is measured, not assumed.
- Cross-channel contact frequency limits are automatically enforced and audited.
- Recommendations appear with reasons; Officer feedback is recorded and used.
- A privacy review (basis of processing, purpose limitations) was completed prior to the pilot.

VISUAL TAKEAWAY

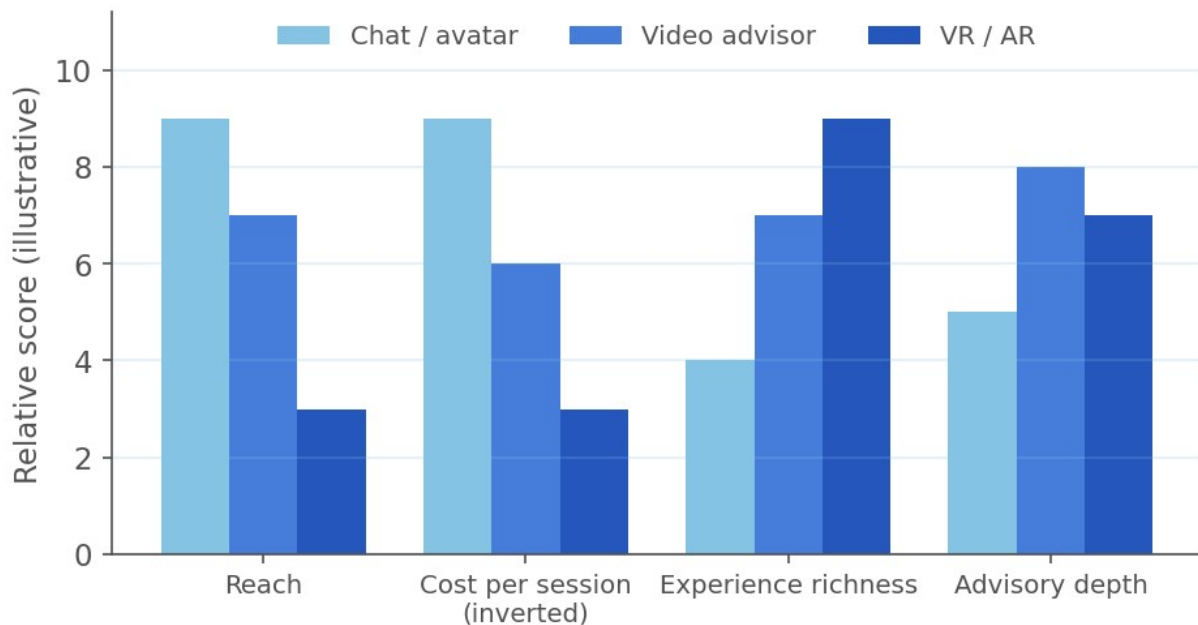
Personalization Is a Closed-Loop Decision System

CHAPTER 6 — Virtual Financial Advice: Financial Advice in a Virtual World



Figure 7. Virtual Financial Advice Journey — Combine goals, suitability, simulation, and human escalation in one auditable flow.

6.1 Business Context



Advice modalities compared (illustrative)

This capability addresses the classic problem of financial advice distribution: quality human advisors are rare and expensive, so only priority clients get them, while regular video calls feel flat for conversations as important as retirement planning or a first mortgage. Virtual Financial Advice presents a middle ground: clients meet advisors — real humans appearing as avatars, or fully AI advisors — in immersive virtual/augmented reality spaces: sitting together in a digital “advising room,” viewing simulated future cash flows as touchable three-dimensional graphs, or walking through a mock-up of a home with a mortgage being calculated.

The business value: expansion of advisor reach (one advisor serving clients across cities without travel), brand differentiation in young and affluent segments, much higher levels of engagement than traditional teleconferencing for complex topics, and a rich engagement footprint for advanced personalization. Strategic honesty is also necessary: the wave of “metaverse” hype has receded from its peak, and the industry lesson is clear — what survives are implementations that solve real problems (remote collaboration for complex decisions, training) and that don't force customers to buy headsets: experiences should be multi-device, accessible from browsers and phones with graceful degradation, with VR as the premium layer, not the sole gateway.

6.2 Technical Workflow

VISUAL TAKEAWAY The system should support judgment, not conceal uncertainty or bypass suitability obligations.

Five technology blocks make up this experience:

Real-time 3D environment. Built on a game engine (Unreal/Unity) or OpenUSD (Universal Scene Description) standards-based 3D collaboration platform that makes it easy to exchange assets between tools — an approach popularized by the NVIDIA Omniverse ecosystem. High-quality rendering (realistic lighting, ray tracing) is run on GPU servers and streamed as pixels to consumer devices (pixel streaming/cloud XR), so even phones and lightweight headsets get workstation-class visuals; interactions are sent back as input. The technical budget is tight: for comfortable VR a high frame rate and low motion-to-photon latency (tens of milliseconds) are required, which means the rendering server should be close to the user (edge/local region).

Avatar and presence. The human advisor is represented by an avatar driven by face-hand tracking from a regular camera (no special tools), with automatic lip-sync of the audio. AI advisors use digital human technology — discussed in depth in Chapter 8 — which is here upgraded to full 3D rendering in shared spaces.

Conversational intelligence and co-pilot. ASR transcribes real-time conversations; LLM + RAG on a product knowledge base serves as an advisory co-pilot (flinging relevant numbers,

documents, and simulations to the virtual “desk” as topics arise) or as a full-fledged AI advisor for mass-market segments with tight compliance guardrails; TTS gives voice to AI advisors. Meeting summaries and follow-up are generated automatically for advisor approval.

Immersive financial visualization. The core value of the experience: a financial simulation engine (portfolio Monte Carlo projections, mortgage amortization, retirement scenarios) whose parameters customers can change directly — slide the “retire at 55 vs 60” slider and see the wealth curve change instantly in three dimensions. In augmented reality mode, the same layer is projected onto the real world: pointing the phone's camera at a property brochure brings up a simulation of the installment above it.

Banking and compliance integration. The session is connected to the bank's CRM/advisory system via API: customer risk profile, actual portfolio, product catalogue; every recommendation that appears passes through the same suitability engine as other channels; and the entire session — audio, transcript, displayed artifacts — is recorded as an audit trail just like a conventional advisory meeting.

6.3 Reference Architecture

1. Client device (browser/phone/headset) ↔ lightweight streaming client.
2. GPU-powered edge render farm: 3D engine instances per session, low-latency video encoder, WebRTC signal server.
3. Session services: advisor-customer matchmaking, scheduling, strong authentication (financial session = full identity verification).
4. Intelligence services: ASR/TTS streaming, LLM+RAG co-pilot, financial simulation engine (stateless, called via API).
5. Core integrations: CRM, advisory/suitability engine, documents and e-signatures for in-session transaction closure.
6. Compliance: encrypted session recording, compliant retention, artifact watermarking, decision logs.

6.4 Success Metrics

Adoption (sessions/month, session completion rate, cross-device), experience (satisfaction scores, motion sickness complaints < threshold, session entry wait times), technical (motion-to-photon latency, stable frame rate, session failures), and business (advisory conversion→transactions vs traditional video channels, AUM per advisor, geographic reach of advisors). Always compare it to the video call baseline — immersion must be proven to add value, not just be novelty.

6.5 Implementation Stages

1. Internal collaboration pilot — use the technology first for advisor training and design meetings (synergy with digital twin Chapter 7) to mature operations without customer risk.
2. Segmented customer pilots — one high-value journey (e.g. mortgage consulting or retirement planning), browser access first, human advisor + AI co-pilot.
3. Expanding the experience — headset support, immersive simulation visualization, AI advisors for heavily gated mass-market segments and seamless escalation to humans.
4. Transaction closing integration — in-session documents and electronic signatures, with full legal review.

6.6 Risks and Mitigation

- Novelty without value — discipline on real problems (complex decisions, distance); measure against video calls; dare to stop features that don't add results.
- Device exclusion — multi-device design from the start; VR optional.
- AI advice errors — mandatory suitability fences in the pipeline, restricted topic coverage, source citations, escalation to humans; for risky products, AI is just the co-pilot.
- Immersive space privacy — facial/motion tracking data is behavioral biometric data: minimize it, don't keep it raw, explain it in consent.
- Session security — strong authentication, end-to-end encryption of streaming, protection against avatar impersonation (verification of the advisor's identity by the system, not by the avatar's appearance).

6.7 Chapter 6 Checklist

- Clear value use cases and measurable video calling baselines.
- The experience runs fine on browser/phone; VR as an upgrade.
- The rendering latency budget is met from the target user's location.
- In-session recommendations pass through the same suitability engine as other channels.
- Sessions are recorded, retained, and auditable; Behavioral biometric data is minimized.
- Proven AI→human escalation path.

VISUAL TAKEAWAY

Advice Earns Trust Through Progressive Guardrails



Discover



Assess



Recommend



Explain



Confirm / escalate

CHAPTER 7 — Digital Twins: Digital Twins of Bank Branches

CHAPTER 07

DIGITAL TWINS: DIGITAL TWINS OF BANK BRANCHES

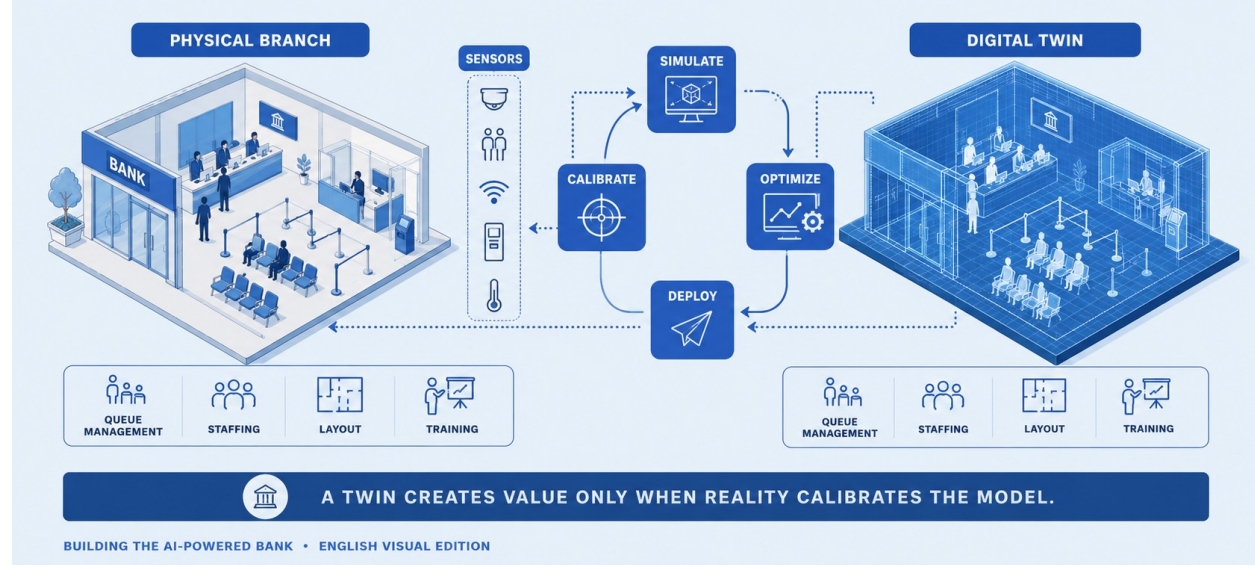


Figure 8. Digital Twin Feedback Loop — Operational data continuously improves simulation, design, training, and branch execution.

7.1 Business Context

A digital twin is a virtual replica of a physical asset or process that is synchronized with real-world data, so that we can observe, simulate, and optimize real systems without touching them. For banks, the object is a network of branches and ATMs: large-value physical assets whose decisions — layout, number of tellers, placement of machines, design of new branches — have long been made by intuition and expensive real-world trials. NVIDIA's vision map calls out its two main uses: training employees and devising new designs — and industry practices add a third: optimizing daily operations.

Example of concrete value: simulating a thousand operational days to find a queue configuration that cuts peak wait times 30% before a single seat is moved; train new tellers to face robbery scenarios or angry customers in VR without risk; testing new format branch designs against real customer traffic before signing tens of billions in renovation contracts; or perform "what-if" branch closures against neighboring branch loads.

7.2 Technical Workflow

VISUAL TAKEAWAY A twin is valuable when its assumptions are calibrated against reality and tied to decisions the bank can change.

3D model layers and semantics. Branches are modeled from architectural drawings (BIM/CAD) or laser scanning/photogrammetry, and assembled in exchange formats like OpenUSD that allow multiple tools (CAD, simulation engines, rendering tools) to work on a single source of truth — patterns that are at the core of platforms like NVIDIA Omniverse. What differentiates a twin from just a mockup: each object has semantics (it's teller counter #2, queue capacity is 8, connected to door sensor X) so that simulation and operational data can "stick" to it.

Live data layer. Twin is fed real data streams: queuing systems, transactions per teller, staff schedules, crowd/door sensors, ATM/CRM machine status, and — when available — metadata from video analytics (number of people per zone, not identities; direct synergy with Chapter 9). This synchronization is what turns mockups into twins: opening a branch twin means looking at the current state.

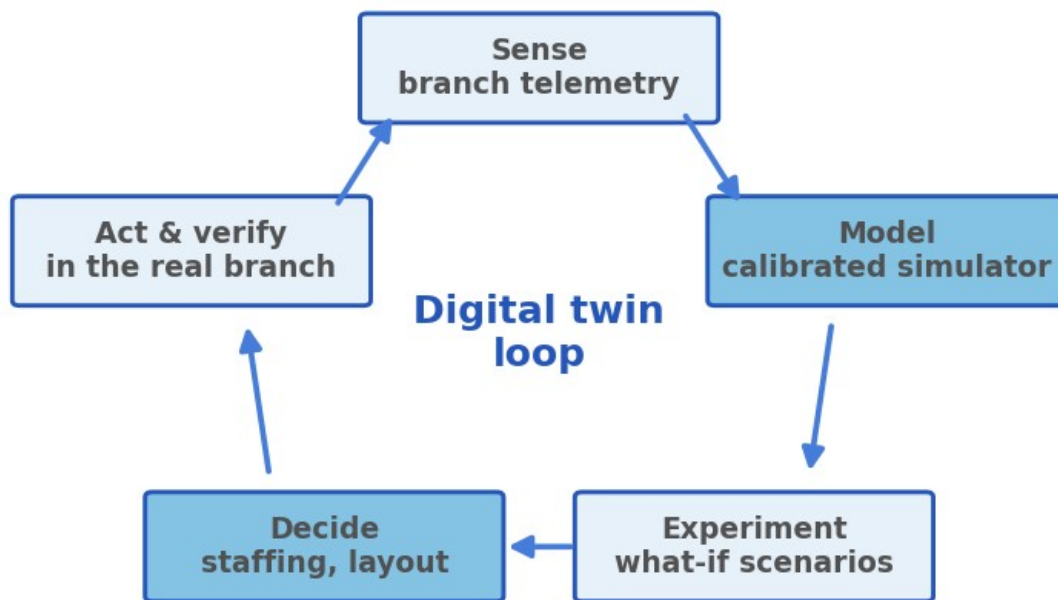
Simulation layer. Three families of engines work on it: (1) discrete-event simulation for queue flow and service processes — customers arrive following the historical arrival distribution, are served following the service time distribution per transaction type; (2) agent-based simulation for individual behavior and spatial movement (walking paths, crowds, evacuation); (3) physical/visual simulations for lighting, camera line of sight, and ergonomics. Simulation parameters are calibrated from historical data and validated: a good twin reproduces past metrics (wait time, utilization) before reliably predicting the future.

Layers of optimization and AI. On top of the simulator, the algorithm searches for the best configuration: combinatorial/metaheuristic optimization for staff scheduling and layout; reinforcement learning that trains in twins (simulation = safe and cheap gym) for dynamic policies such as adaptive counter open-close; and demand prediction models (hourly customer arrivals, ATM cash requirements) that feed the scenarios. This “train on twin, apply in the real world” pattern — sim-to-real — is one of the most powerful uses of twin, including for training and testing Chapter 9 computer vision models with synthetic data: twin can render thousands of scene variations (lighting, crowds, rare events like a bag being left behind) complete with perfect labels, something that would be impossible to gather from real CCTV.

Layers of experience. Twin is accessed via an interactive 3D dashboard for operations managers, a collaborative design mode for architects (several people change the layout together), and a VR mode for employee training: service, compliance, and emergency

scenarios with conversational AI-driven customer avatars (Chapter 8 synergy), complete with automatic scoring of participant responses.

7.3 Reference Architecture



The digital twin loop: sense, model, experiment, decide, verify

1. Source: BIM/CAD + scan → USD model; IoT/queue/transaction → event bus.
2. Twin platforms: USD scene storage, object-data synchronization service, asset catalog (furniture, machines, avatars).
3. GPU compute farm: collaborative rendering and VR; parallel simulation execution (thousands of scenario replicas); synthetic data generation.
4. AI services: demand prediction, optimization engine, RL environment, training evaluator.
5. Consumer: operational dashboard, design studio, VR training portal, export of results (schedule recommendations to HR system, configuration to facility management).
6. Governance: scene versioning ("which designs are approved" audits), per-branch access controls, and data policies (twin uses aggregated/anonymous data — not individual customer tracking).

7.4 Success Metrics

Twin validity (difference in simulated vs. actual metrics < threshold, e.g. 10%), operational impact (decreased lead time, increased staff utilization, schedule savings), design impact (avoided change costs, design cycle speed), training (competency scores, knowledge retention, new hire time-to-work vs old methods), and adoption (frequency of decisions citing twin results).

7.5 Implementation Stages

1. Pilot single-branch twin — 3D model + queue/transaction data; simulator calibration and validation; one real decision (staff schedule) is taken from the twin.
2. VR training module — 3–5 priority scenarios with assessments; compare learning outcomes vs conventional classes.
3. Design studio — standard branch component library; test new format designs with simulated traffic before development.
4. Network scale — twin creation templatization (scan-to-twin pipelines), cross-branch portfolio dashboards, RL for operational policies, synthetic data factories for vision models.

7.6 Risks and Mitigation

- Twins are beautiful but false — without calibration-validation, simulations mislead decisions; make validation a routine ritual and report confidence intervals, not single numbers.
- Modeling costs balloon — from the level of detail the decision requires (queues don't need marble textures); standardize components; automate scan-to-twin.
- Stale data — twins without synchronization become mockups; set data freshness and pipeline monitoring SLAs.
- Privacy — use aggregate data; If video metadata is used, ensure it is anonymous and complies with Chapter 9 privacy review.
- Low adoption — involve branch managers from design; start from decisions that are painful for them (schedules, queues), not from beautiful demos.

7.7 Chapter 7 Checklist

- 3D model source and exchange format (USD) defined; defined object semantics.
- The simulator is calibrated and validated against historical data.
- At least one real operational decision per quarter is taken based on the twin.
- The VR training module has an assessment rubric and compares old methods.
- Twin data policy: aggregated, anonymous, clearly retained.
- Synthetic data pipelines connect to vision model requirements (Chapter 9/13).

CHAPTER 8 — Avatar Assisted Banking: Digital Humans Serving Customers



Figure 9. Banking Avatar and RAG Pipeline — A controlled conversational system separates language generation from trusted banking actions.

8.1 Business Context

The old generation of menu-based chatbots left a bad legacy: customers got stuck in a “type 1 for balance” loop, got frustrated, then kept calling the call center. Avatar Assisted Banking is a leap of two levels at once: from menus to conversational AI that truly understands natural language, and from text to digital humans — avatars with faces, voices, and expressions that talk to customers at branch kiosks, next-generation ATM screens, mobile apps, or websites. The psychology is real: face and voice increase trust, comfort for customers who are less familiar with typing, and accessibility for the elderly and those with limited digital literacy.

The business value rests in the cost and experience quadrant: high containment rate (percentage of interactions completed without humans) cuts the cost per contact from tens of thousands of rupiah (human agents) to hundreds of rupiah; service becomes 24×7 without queuing; consistency of answers increases; and human agents are freed up for complex, high-value cases. In branches, avatars absorb routine transactions (product information, card status, form guides) so teller lines shorten — in direct synergy with NBA Chapter 5.

8.2 How It Works: Conversation Pipelines

VISUAL TAKEAWAY Retrieval grounding and tool permissions reduce hallucination risk, but high-impact actions still require explicit confirmation.

A single turn of spoken conversation goes through the following chain, with a total latency budget ideally under a second or two to make it feel natural:

1. **Audio capture** — microphone array with noise suppression and echo cancellation (noisy branch environment); voice activity detection intercepts speech.
2. **ASR streaming** — staged transcription while the customer is still talking (rather than waiting for it to finish), with a model fine-tuned for conversational Indonesian, a mix of English terms and banking vocabulary; word boosting for product names.
3. **Comprehension and dialogue** — two architectures apply: (a) classic NLU (intent classification + entity extraction + dialogue state machine) — deterministic, easy to audit, suitable for strictly transactional flows; (b) LLM + RAG — much more flexible understanding, answers from the official knowledge base with citations, and function calling to call banking APIs (check balance, block card) in a structured manner. Mature practices combine both: LLM as understander and orchestrator, critical flow (transaction, verification) remains through a rigid state machine.
4. **Guardrails** — before and after LLM: topic filters (still in the banking realm), prompt injection detection, personal data filters on output, checking the suitability of answers to source documents (groundedness), and escalation policies. The ecosystem provides ready-made frameworks (e.g. NeMo Guardrails-style guardrails patterns) — but the content policies still have to be written by the bank.
5. **Expressive TTS** — natural Indonesian voice with intonation control; consistent brand voice across channels.
6. **Digital human animation** — TTS audio animated lip-sync and facial expressions (audio-to-face technology), plus gestures and gaze; 2D/3D rendering at the branch edge or streamed from GPU servers — a family of technologies popularized by human digital platforms like NVIDIA ACE/Tokio. For text channels (app/web), the same pipeline runs without stages 1-2 and 5-6.

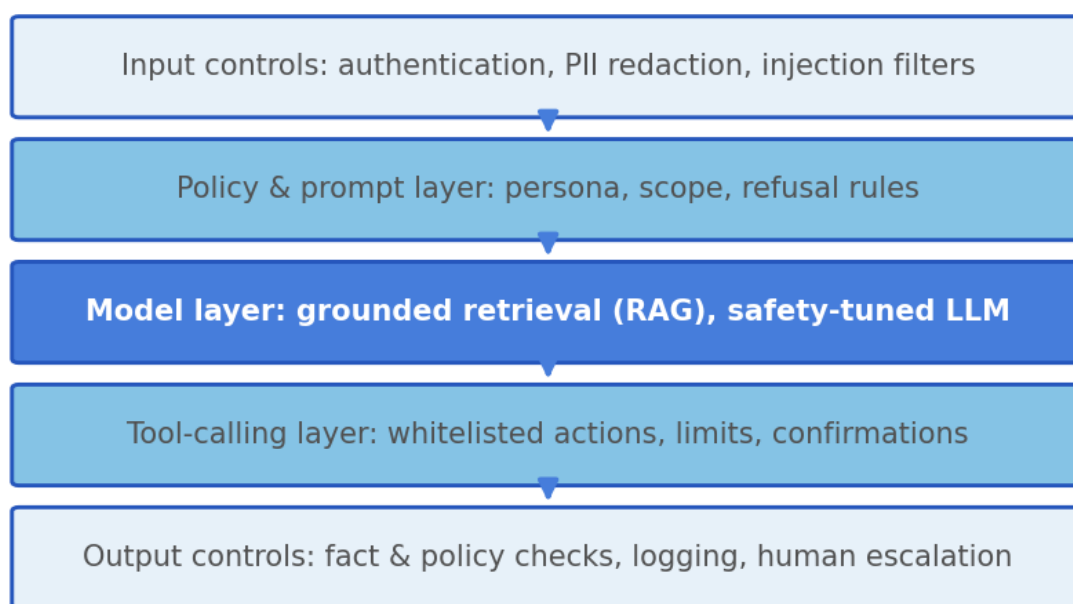
8.3 The Brains Behind It: RAG and Action Orchestration

The avatar knowledge base is built from authoritative sources: product documents, FAQs, procedures, terms, and fee policies — chunked with structure, embedded, and indexed in a vector database; hybrid retrieval (semantics + keywords) plus re-ranking provides the best context to the LLM; answers must cite sources, and out-of-scope questions are answered honestly "I don't have that information" plus an offer of escalation. Knowledge updating becomes a matter of document pipelines (publish → index in minutes), not model retraining.

For actions, the avatar calls the API through a defined scheme (function calling) with multi-layered security principles: the customer's identity is verified by the bank's authentication system (not by

the LLM), each function has a whitelist of parameters and value limits, risky actions (transfers, data changes) require explicit re-authentication and confirmation, and all calls are logged. LLMs never hold credentials; it just asks, the system decides.

8.4 Reference Architecture



Layered guardrail architecture around the banking avatar

1. Interaction points: branch kiosk / ATM screen / app / web — light client (two-way audio-video streaming).
2. Branch edge (optional): small GPU server for ASR/TTS/local rendering — maintains latency and operations when the WAN is compromised; synchronize policies from the center.
3. Conversation core (central/private cloud): session gateway, ASR/TTS service, dialogue orchestrator, LLM inference (optimized for throughput with in-flight batching), RAG stack (embedding, vector DB, re-ranker), guardrails.
4. Integration: banking API via gateway with strict authorization; authentication system; handover to a human agent (the context of the conversation changes — don't force the customer to repeat the story).
5. Observability: conversation logs (with personal data redaction), per-stage latency metrics, intent analytics, human review queue for conversation samples.

8.5 Success Metrics

Category	Metric	Reasonable Initial Target
Effectiveness	Containment rate	40–60%, rising over time
Quality	Correct resolution (sample audit)	≥ 90% within audited coverage
Honesty	Hallucination rate / groundedness	Near-zero hallucinations in testing
Experience	Post-interaction CSAT, task completion	≥ legacy text channel
Technical	Speech-to-response latency p95	≤ 1.5–2 seconds
Security	Injection / unauthorized-action incidents	0; red-team tested
Escalation	Seamless handover with full context	≥ 95%

Also measure honest deflection: make sure containment goes up because problems are solved, not because customers give up — check with a 48-hour recontact rate.

8.6 Implementation Stages

1. **Text phase** — LLM+RAG assistant on app/web for light information & services; build knowledge bases, guardrails, automated evaluations (golden questions), and human review processes.
2. **Voice phase** — add ASR/TTS in IVR and apps; fine-tune ASR on real audio; chase latency budgets.
3. **Branch avatar phase** — digital human kiosks in pilot branches; queue-to-stall design; measure the absorption of routine transactions.
4. **Action phase** — function calling for low → medium risk transactions with layered authentication; expand coverage based on failure data.
5. **Optimization phase** — personalization (customer context-aware via store features, permission), multilingual/regional, and NBA integration (avatars convey relevant offers according to Chapter 5 rules).

8.7 Risks and Mitigation

- Financial answer hallucination — mandatory RAG, mandatory citations, restricted coverage, automated groundedness evaluation in CI, daily samples human reviewed; answer numbers (fees, interest) only from the API/document, never from the "memory" of the model.
- Prompt injection & abuse — strict separation of instructions vs data, input/output filters, function and value boundaries, re-authentication for actions, periodic red teaming (Chapter 16).
- Personal data leaks — redaction of PII in logs, short retention periods for audio, prohibition of customer data from going to third party prompts without agreement and control.
- Uncanny experience / accessibility — test avatar designs with real customers; always provide alternative paths (text, people); do not deny access to services for those who reject AI.
- Model vendor dependencies — model layer abstraction (easy to replace LLM), periodic evaluation of multiple candidates, full ownership of knowledge base and logs.
- Communication compliance — all offer and disclaimer templates are compliance approved; the avatar reveals itself as a digital assistant (AI transparency).

8.8 Chapter 8 Checklist

- Official, curated knowledge base with a fast update pipeline.
- Active input-output guardrails; red team injection passed before release.
- Latency budget per stage is measured (ASR/LLM/TTS/render) and met.
- API actions: authentication by bank system, explicit confirmation, complete logs.
- Human handovers bring context; recontact rates are monitored.
- Automatic evaluation (golden set) runs on CI; Conversation samples are reviewed daily.

VISUAL TAKEAWAY

The Avatar Is the Interface - Grounding Is the Product



CHAPTER 9 — Intelligent Video Analytics for Threat Detection in Branches

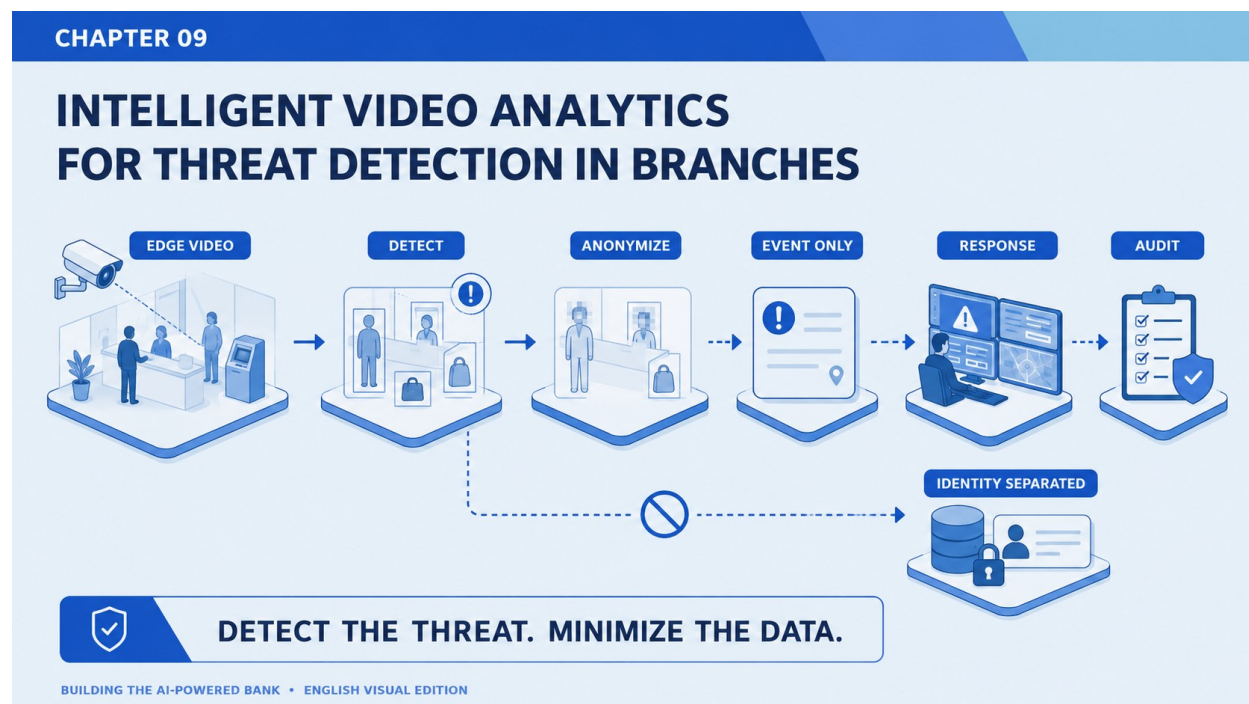


Figure 10. Privacy-Aware Video Analytics — Process the minimum data required and separate detection from identity whenever possible.

9.1 Business Context

Bank branches are physical risk environments: robberies, violence against staff, theft, suspicious objects, even medical emergencies. The traditional security model — security guards plus a wall of human-watched CCTV monitors — has biological limits: human factors studies show operators' attention to multiple screens plummets in just tens of minutes. Intelligent Video Analytics (IVA) flips that model on its head: machines watch all the cameras every second and raise events that need attention to humans — from “humans looking for events” to “humans looking for events.” According to the vision map: IVA analyzes the behavior of individuals within the branch to identify suspicious activity and abandoned objects.

Its business value is in the risk and cost quadrant: minute-to-second incident response, prevention (early intervention in escalations), quality evidence for investigations, workplace safety compliance, and security team efficiency (one command center overseeing hundreds of branches). Operational bonus: the same metadata (person count, density, queues) fuels the digital twin Chapter 7 and service analytics.

9.2 How It Works: AI Video Pipeline

VISUAL TAKEAWAY Privacy by design starts with purpose limitation, short retention, controlled access, and auditable response workflows.

The default software architecture for video AI is a streaming pipeline (standardized pattern of frameworks like NVIDIA DeepStream on top of GStreamer): hardware-accelerated decoding of multiple RTSP streams → batch pre-processing → chained inference → tracking → event logic → metadata output. The main stages:

1. **Per-frame object detection:** people, bags/suitcases, and other relevant classes — optimized single-stage detection model (via TensorRT, INT8/FP16 precision) allows dozens of channels per GPU.
2. **Multi-object tracking (MOT):** tying inter-frame detection into IDed trajectories (SORT/DeepSORT/ByteTrack family algorithms with re-identification appearance features) — the foundation of all behavioral analytics: without trajectories, there is no “how long is this person here”.
3. **Behavioral analytics on the track:**
 - **Loitering** — staying in a certain zone exceeding a time threshold (e.g. ATM area > X minutes without a transaction).
 - **Landless object** — a static object appears, its owner (the track that left it) moves away beyond the distance-time threshold.
 - **Zone & direction intrusion** — entering restricted areas, going against the flow in one-way lanes, entering staff-only areas.
 - **Action recognition** — action models (3D-CNN/video transformer or sequence pose classification from pose estimation) for critical classes: falling (medical emergency), panic running, fighting/violence, raising an object, hands above head (indication of robbery). These rare classes are trained with the help of synthetic data from digital twins (Chapter 7) and augmentation, as real recordings are almost non-existent.
 - **Statistical anomalies** — the model learns each camera's “normality” (hourly patterns of density and motion) and flags large deviations — a safety net for previously undefined events.
4. **Fusion and event logic** — the rules engine combines signals (detection + zone + time + operating hours context + door alarm signal) into severity-rated events, with hysteresis and de-duplication to not overwhelm operators.
5. **Outputs** — events + evidence clips + structured metadata are published to the command center (VMS/PSIM), branch officer applications, and the enterprise event bus.

9.3 Privacy Principles by Design

Branch IVAs can be established without individual identification — and preferably so, unless there is a strong legal basis for a particular case:

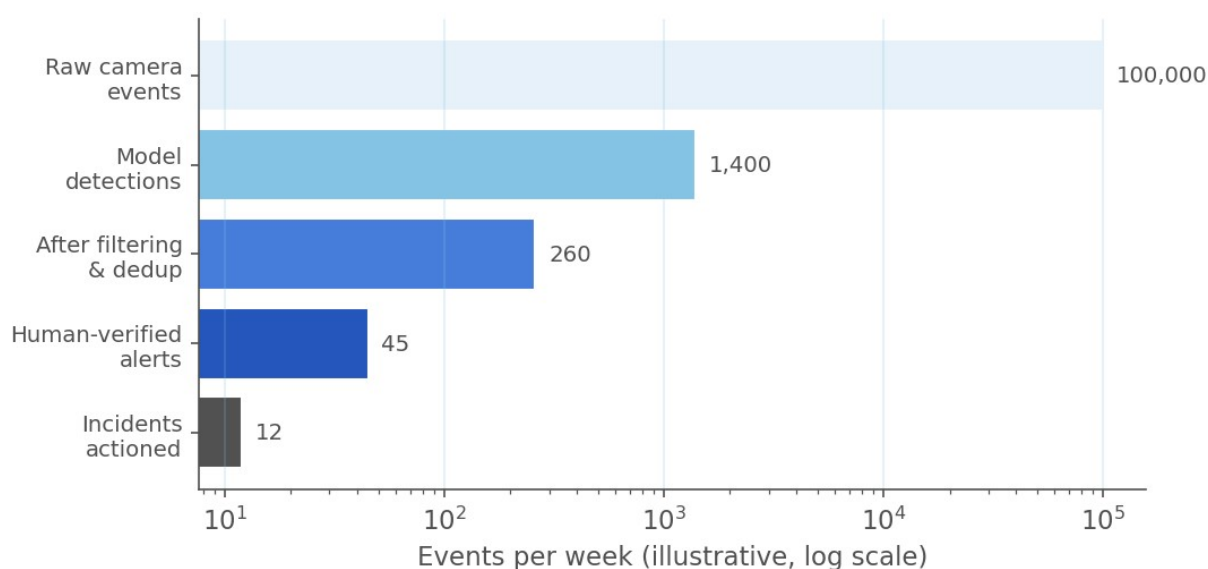
- Behavior and object-based analytics, not faces; there is no identity matching in normal operation.

- **Minimization:** what leaves the branch is metadata and event clips, not continuous raw video; short raw retention at the edge.
- **Automatic redaction (face blur) for secondary uses:** model training, service analytics, sharing to third parties.
- **Transparency:** monitored area notification signage, written policies, data protection impact assessment (DPIA) before title, role-based record access with logs.
- The ban on disproportionate functions (e.g. customer demographic profiles from cameras) is enforced as a technical policy, not just an advisory.

9.4 Reference Architecture

1. IP cameras (RTSP, placement designed from line-of-sight studies — can be simulated in a digital twin).
2. **Edge server GPU per branch: decode-detect-trace-analytics pipeline; local record buffer; standalone operation when WAN goes down (store-and-forward event).**
3. Central: event aggregation, command center (event-based video wall, not camera wall), cross-branch correlation, chain-custodial evidence storage (hashes, timestamps, access logs).
4. Model platform: vision model registry, training pipeline (labeled real data + synthetics), per-camera evaluation, OTA model distribution to edge fleets with incremental releases and rollbacks.
5. Integration: alarm/PSIM, incident ticket system, officer notification, and (aggregated metadata) to digital twin and operational analytics.

9.5 Data, Models, and Evaluation



From raw events to actioned incidents: the alert funnel (illustrative)

Data requirements: representative footage of each camera/lighting type for detection fine-tuning; zone annotation per camera (drawn once, diversion); a dataset of labeled events from running operations (every alert the operator reviews = new label — build this loop from day one); synthetic data for rare classes. The evaluation uses two layers of metrics: per-model (mAP detection, MOTA/IDF1 tracking, action accuracy per class) and — which determines the fate of the project — per-operational event: alert precision (of alerts raised, how many were correct), event recall (of real incidents, how many were detected — tested with drills/recreated events), detection time, and alert load per operator per hour. A system with low precision will be abandoned by operators in a matter of weeks (alert fatigue) — it is better to have a narrow but reliable event coverage, then expand it.

9.6 Implementation Stages

1. Pilot 1–3 branches, 3 clear-impact use cases (special area intrusion, no man's land, crash/emergency) — calibrate thresholds per camera, train operators, measure precision/recall with test scenarios.
2. Operational hardening — response SOPs per incident type, command center integration, chain of custody of evidence, final DPIA.
3. Class expansion — loitering, violence/robbery indicators (with synthetic data), anomalies; operator feedback based tuning.
4. Fleet scale — edge provisioning automation, health monitoring (camera outages, image quality drift), OTA model distribution, cross-branch risk portfolio dashboards.

9.7 Risks and Mitigation

- Alert fatigue — pursue precision first; aggregation/de-dup; adaptive threshold per camera; review weekly load alerts.
- Lane performance between conditions — evaluation per camera/lighting/segment; augmentation and additional data for weak points; don't believe the average.
- Adversarial & sabotage attacks — closed/shifted/looped camera detection; edge device hardening; flow integrity (signing).
- Privacy & mission breaches widen — technical policies limit functionality; usage audits; every new function must be reviewed.
- Connection dependency — edge-first: local detection and alarms stay on when the WAN is down.

9.8 Chapter 9 Checklist

- DPIA completed; the principles of non-identification and minimization are technically enforced.
- The edge pipeline meets channel/GPU targets with an optimized model.
- Precision alerts and operator alert load are monitored; label loop of running operator overview.
- Critical event classes are tested against planned scenarios (not just offline metrics).

- Chain-of-custody proof: hash, timestamp, access log.
- OTA model distribution with gradual releases and tested rollbacks.

VISUAL TAKEAWAY

From Camera Pixels to an Auditable Case

EVENT METADATA + TIMESTAMPS + MODEL VERSION + HUMAN DECISION

CHAPTER 10 — Data Visualization for Risk Management and Predictive Analytics

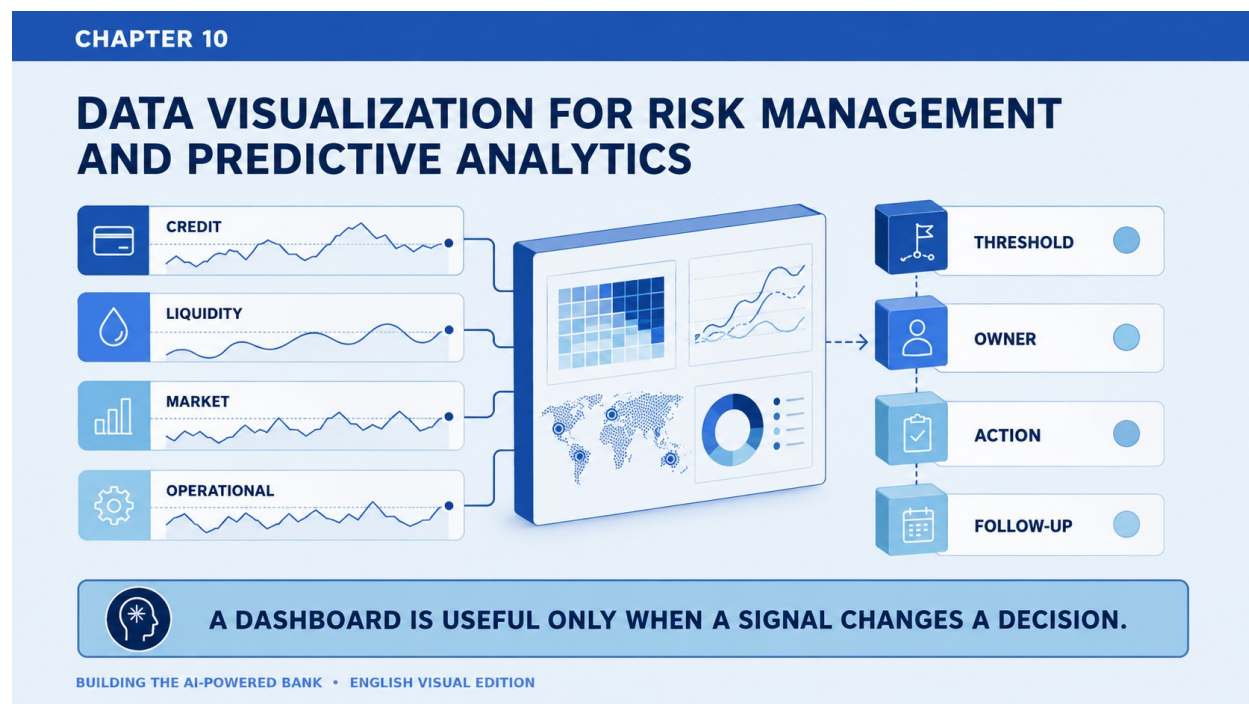


Figure 11. Risk Analytics Signal Chart — Illustrative emphasis by signal domain—not a performance benchmark.

10.1 Business Context

In the vision map, this capability is depicted simply — screens full of graphics in the workspace — but its role is fundamental: data visualization improves risk management and drives predictive analytics. Banks are in the business of managing risk with information; the decisions of the credit committee, treasury, and board of directors are only as good as their ability to discern. The classic problem: risk reports in the form of static monthly tables that are late, fragmented between silos, and cannot be answered with follow-up questions ("why did this segment's NPL increase?") without waiting for the analyst a week.

Its modern capabilities change three things at once: freshness (from monthly to daily/real-time for fast-moving risks like liquidity, markets, and fraud), interactive depth (drill-down from aggregates to transactions in seconds), and viewing direction (from the rearview mirror — what happened — to the windshield: predictive analytics that project what will happen, embedded directly in the same dashboard).

10.2 Substance: What is Visualized

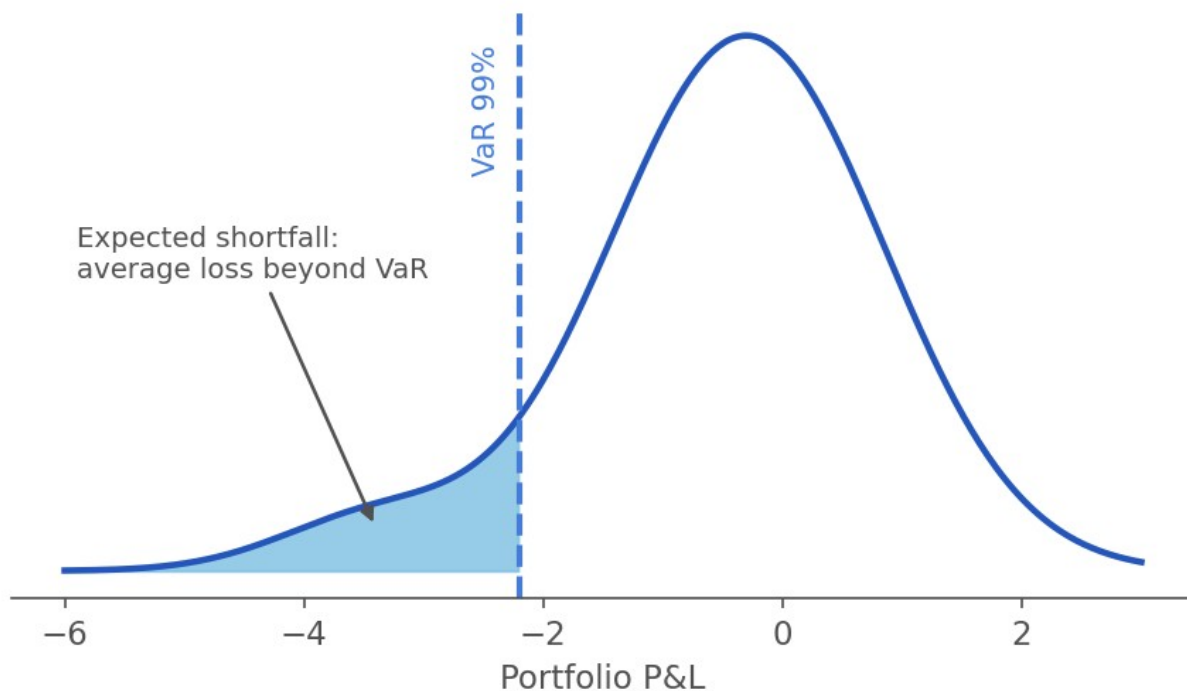
VISUAL TAKEAWAY Good dashboards connect leading indicators to decisions, thresholds, owners, and follow-up actions.

Risk framework that forms the contents of the dashboard:

- Credit risk — portfolio quality (NPL, coverage), collectibility migration (transition matrix), vintage analysis per disbursement cohort, concentration (sector, debtor group, geography), expected loss parameters (PD, LGD, EAD) and their movements, macro scenario stress test results.
- Market & liquidity risk — Value at Risk and sensitivity, repricing gap, liquidity ratio (LCR/NSFR), cash flow projections, behavior of third party funds.
- Operational risk & fraud — incident and loss trends, process heat maps, real-time fraud metrics (attack rate per channel, prevented losses — synergy Chapter 11), model health (drift, performance — synergy Chapter 15).
- Compliance — AML alert queues and their lifespan, reporting SLAs, audit findings.

The design principle: each number has an owner, a single definition (semantic layer), and an action threshold; a dashboard that doesn't change anyone's decisions is decoration.

10.3 Technical Workflow



Loss distribution, value at risk, and expected shortfall

Data and semantic layers. The dashboard attaches to lakehouse gold (Chapter 4) through a semantic layer — centralized metric definition (one NPL formula for the entire bank) — with a real-time path from stream processing to fast metrics. This eliminates the “three meetings, three different NPL numbers” disease.

Computing acceleration. Interactive analytics over billions of rows demands speed; This is where GPU-accelerated analytics comes in: the RAPIDS ecosystem (cuDF for dataframes, cuML for ML, Spark acceleration) and GPU-accelerated visualization engines enable filter-aggregation-render over hundreds of millions of points in sub-seconds — changing analyst work patterns from “query, wait, drink coffee” to fluid exploration. Heavy simulations (Monte Carlo VaR, stress tests of thousands of scenarios, XVA) overnight on the CPU are down to minutes on the GPU — enabling intraday risk that was once impossible.

Embedded predictive layer. Forecasting models (time series: from classical statistical methods to temporal transformers) project NPLs, cash flows and volumes; the early warning model scores debtor/risk segments upwards; and explainability (SHAP value) is displayed next to the prediction: not just “NPL segment X will rise”, but “driven by weakening sector Y and increase in limit utilization”. Predictions without explanation won't be trusted by committees — and shouldn't be.

Consumption layer. Per-person interactive dashboards (directors: concise + scenarios; risk managers: drill-down; analysts: notebooks connected to the same data), threshold and anomaly-based alerting (users don't have to open the dashboard to know there's a problem), and — the latest development — natural language interfaces: LLM on top of the semantic layer translates “show mortgage NPL trends per region 12 months” into validated queries, with guardrails: LLM only builds queries on official metrics, doesn't make up numbers; results always show the query/definition used.

10.4 Reference Architecture

1. Lakehouse gold + stream metrics → semantic layer (metric definition, row/column access control).
2. Query engine (warehouse/lakehouse engine; GPU acceleration for heavy interactive loads) + risk compute cluster (GPU simulation).
3. Predictive model services via inference server; SHAP annotations are calculated and stored alongside predictions.
4. Serving layer: BI platform, custom dashboard application, metrics API for other applications, NL-to-query gateway.
5. Governance: metrics catalog with owner and lineage, official dashboard certification, access traces for sensitive data.

10.5 Success Metrics

Freshness (data lag per metric vs target SLA), speed (sub-second p95 interactive query latency for exploration), trust (single source of numbers — loss of manual reconciliation between units), adoption (active users per persona, meeting decisions citing dashboards), and predictive impact (forecast vs actual accuracy, lead time early warning before worsening collectibility, loss avoided by early action).

10.6 Implementation Stages

1. One domain, one truth — choose credit risk; build a semantic layer + 10–15 certified core metrics; retire counter report.
2. Interactivity and freshness — transaction drill-down, daily/real-time track for fast metrics; GPU acceleration when volume demands.
3. Embedded predictive — forecast + early warning + SHAP in the same dashboard; clear alert follow-up process.
4. Domain expansion & NL interface — markets, liquidity, operations; Natural language interface on top of a semantic layer with fences.

10.7 Risks and Mitigation

- Decisionless dashboards — each dashboard has an owner, business question, and action threshold; quarterly adoption audits, retire the dead.
- The multiple-versions-of-truth crisis — mandatory semantic layer; certified "official" dashboard; Metric definitions change through governance.
- Misleading predictions — show uncertainty and explanation intervals; monitor forecast accuracy; don't automate big decisions from forecasts without humans.
- Data leakage via BI — control row/column access at the semantic layer, not on each dashboard; access log; masking for a non-authoritative role.
- NL interface fabricates — limit to official metrics, display queries, test question-answer regression.

10.8 Chapter 10 Checklist

- Semantic layer with certified, owner, and lineage metrics.
- Freshness SLAs per defined and monitored metrics.
- Interactive exploration meets latency targets at full volume.
- Predictions appear with confidence and explanation intervals (SHAP).
- Threshold/anomaly based alerts are connected to a follow-up process.
- Access control and audit trail at the semantic layer.

CHAPTER 11 — Fraud Alert: Real-Time Fraud Detection with Deep Learning



Figure 12. Real-Time Fraud Decision Path — Milliseconds matter, but speed cannot replace resilience, observability, and customer protection.

11.1 Business Context

Fraud is an endless arms race. On one side stand organized, well-capitalized, and adaptive syndicates — they test loopholes, buy stolen data, and switch gears in a matter of days. On the other hand, banks must decide in tens of milliseconds, millions of times a day, which transactions to release and which to hold — with two equally grievous costs of error: false negatives mean immediate losses and an erosion of trust; A false positive means a legitimate customer was denied their card in front of the cashier, embarrassed, and possibly changing banks. The vision map summarizes the promise of this capability: deep learning improves fraud detection accuracy and enables real-time notification.

The threat landscape is wide and must be mapped because each mode requires different signals: card fraud (card-not-present, skimming/counterfeit, lost-stolen card), account takeover (stolen credentials, SIM swap, malware), victim-authorization payment fraud (APP fraud: victims are tricked into making their own transfers — social engineering, "fraud under the guise of gifts/officials/romance"), credit application fraud (fake identities and synthetic identities — a combination of real data and fabricated fabrications). for years), internal fraud, merchant fraud, and money mule — intermediary accounts that are the lifeblood of laundering the proceeds of crime. A

good system does not detect “fraud” as a single class, but rather a portfolio of modes with layered models and controls.

11.2 Why Deep Learning (and When Not)

First-generation systems were rules-based: transparent and quick to create, but fragile — cheaters learned the threshold, and the explosion of rules became untenable. The second generation, classic ML (especially gradient boosting over engineered features), improves accuracy significantly and remains a very strong baseline. Deep learning adds three powers that were difficult to achieve before:

1. **Sequence learning** — transformers/RNNs over sequences of customer events (transaction, login, change profile, add recipient) capture the dynamics that are lost in aggregate features: the typical account takeover pattern is a sequence (login new device → change number → add recipient → multiple transfers), not one odd transaction.
2. **Relationship learning** — Graph Neural Network over a graph of entities (accounts, cards, devices, IPs, addresses, merchants) detects a syndicate structure: devices used by 40 accounts, rotating transfer rings, credit application clusters sharing attributes. A node's GNN score “inherits” the suspicions of its neighbors — exactly how senior analysts think.
3. **Novelty detection** — autoencoders and density models flag previously unknown patterns — defense against new modes before the first label arrives.

Winning architectures in production are almost always layered ensembles: hard rules for certainty (blacklists, regulatory limits), gradient boosting for tabular features, sequence models for behavior, GNNs for relations (often calculated near-real-time/periodically and presented as features), and anomaly scores — combined meta-models or decision policies. The principle: each layer catches crimes that escape other layers.

11.3 Feature and Data Engineering

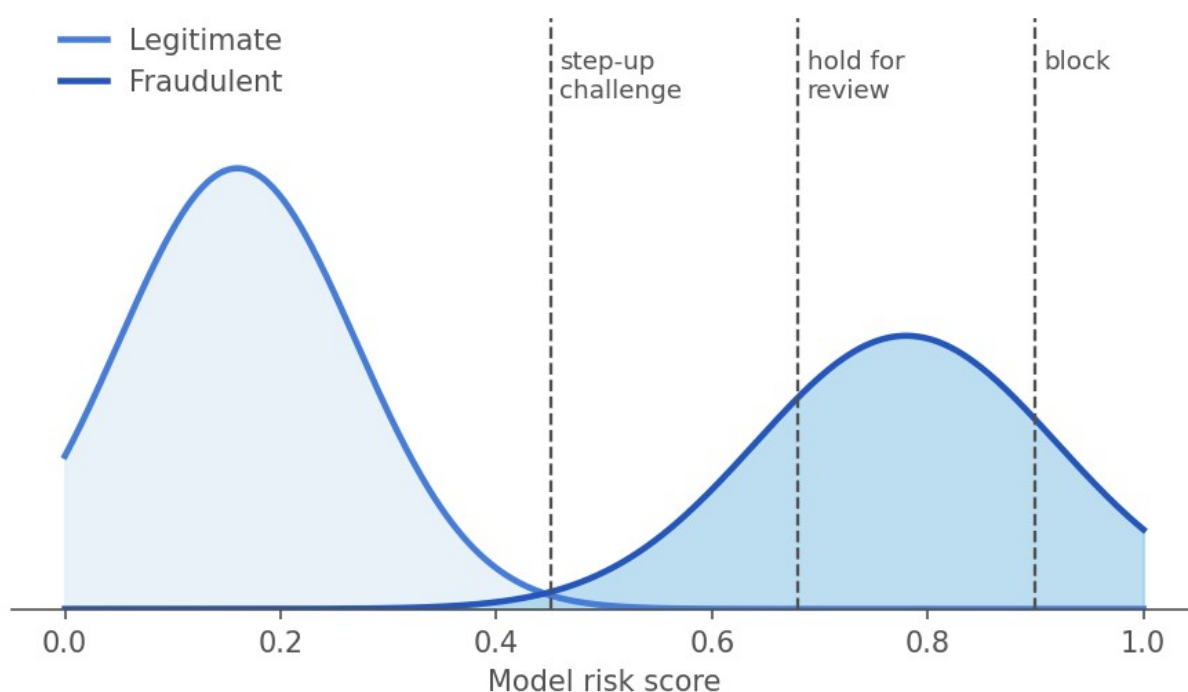
The fuel of the fraud model is the behavioral features and speed (velocity) that the stream computes (Chapter 4):

- **Velocity:** number and value of transactions per 1 minute / 1 hour / 24 hours / 7 days per card, account, device, IP, merchant.
- **Habitual profile:** deviation from customer baseline (time, location, merchant, typical amount — modeled as a distribution, not an average).
- **Device and session signals:** device fingerprint, IP/ASN reputation, impossible-to-move (two unreachable locations in a time interval), behavioral biometrics (typing rhythm, touch patterns — a strong differentiator of real humans vs. bots/fraudsters, with strict privacy requirements).
- **Graph features:** node degree, community density, distance to poorly labeled entities, GNN score.

- External context: running mode reports, list of interbank consortia, SIM-swap signals from operators (when available through official cooperation).

Two data disciplines determine life-or-death: point-in-time correctness — training features must be exactly as available at the time of the event (feature store with point-in-time join; future leaks produce false champion models); and label cleanliness — fraud labels come too late (chargebacks take weeks), some victims don't report them, and definitions must be consistent (third-party fraud vs. dispute vs. first-party abuse are separated). Manage maturity window labels explicitly when building datasets.

11.4 Addressing Extreme Class Imbalance



Score distributions and graduated action thresholds

With fraud prevalence often $< 0.1\%$, the standard practice is: train with class or focal loss weighting (careful negative undersampling; synthetic oversampling like SMOTE is rarely helpful on real transaction data and can harm calibration), evaluate with PR-AUC and recall on alert budgets (e.g. recall@0.5% alert rate) instead of accuracy/ROC alone, recalibrate post-training scores, and set thresholds of the business cost curve: value of prevented loss vs cost of analyst reviews vs cost of customer attrition — The optimal threshold varies per segment and per channel, and should be reviewed periodically as the distribution shifts.

11.5 Real-Time Architecture: Decisive Milliseconds

VISUAL TAKEAWAY The fraud platform should degrade gracefully and preserve safe decisioning when a feature, model, or dependency fails.

Inline decision path on authorization:

1. Switching/payment gateway calls for decision service fraud (typical total budget 20–100 ms depending on scheme).
2. The decision service pulls features online (feature store, single digit ms), runs the ensemble on optimized GPU/CPU inference servers (quantization, micro-batching; sequence models use per-customer state cache), applies per-segment thresholding policies.
3. Decision: release / challenge (step-up: OTP, biometrics, in-app confirmation) / hold. The challenge option is an important weapon: shifting the trade-off from “deny vs release” to fast, customer-friendly verification.
4. **Real-time notifications — as the use case promises: instant push/SMS/in-app for held or suspicious transactions, with “it's me / not me” buttons — customer responses become the cheapest and fastest labels that flow straight back to the system (and cancel the hold in seconds if valid).**
5. Parallel asynchronous pipelines: deep rescoring (full GNN, heavy models), cross-transaction correlation, and case management queues for analysts with auto-collected evidence and score explanations (boosting features via SHAP) — speed up investigations while meeting audit needs.

The ecosystem provides GPU-accelerated blueprints for such pipelines (cybersecurity/fraud frameworks like NVIDIA Morpheus for streaming inference, and large-scale GNN libraries) — but its logical architecture, as outlined above, is vendor-neutral.

11.6 Countering Drift and Adaptive Adversaries

Fraud is the domain of permanent and adversarial drift: an adversary actively attacks your model. Sustainable defense practices:

- **Daily monitoring:** score distribution, alert rate, analyst review precision, and recall-proxy (fraud detected late) per segment/channel; alarm on shift.
- **Scheduled + triggered retraining** (Chapter 15) with mature label window; champion-challenger and shadow deployment before promotion.
- **Three-source feedback:** analyst decisions, customer responses to notifications, and chargebacks — all come back as labels with different reliability weights.
- **Red team internal fraud:** a team that tries to penetrate its own system with new methods; findings become training data.
- **Collaboration:** sharing indicators between banks/consortiums and with law enforcement according to legal corridors — syndicates attack many banks at once; collective defense transformed their economy.

- Beware of adaptive attacks: threshold probing (small repeated transactions to map boundaries), behavioral pollution (building false “habits” before attacking), and manipulation of features that the attacker can control — the design of features that are expensive to fake (graph relations, long histories) is more robust than surface features.

11.7 Success Metrics

Metric	Practical Definition
Fraud loss rate	Fraud losses / transaction volume (bps)
Recall @ alert budget	% value of fraud caught in X% alerts
Precision alerts	Percentage of alerts confirmed as fraud
False positive customer rate	% legitimate transactions declined/challenged
Detection time	Median seconds from event to alert
Case resolution time	Hours from alert to analyst decision
Loss prevented	Held value confirmed as fraud

Always report metric pairs (e.g. recall at fixed FPR) and their trends per mode — single numbers hide trade-offs.

11.8 Implementation Stages

1. **Foundation** — feature store streaming for core velocity; baseline gradient boosting replaces/accompanies rules; case management by recording disciplinary labels.
2. **Full real-time** — inline decision service meets latency budget; step-up challenge; two-way customer notification.
3. **Model depth** — behavior sequence model; entity graph + GNN (start batch, then near-real-time); layer anomalies.
4. **Adversarial maturity** — automatic drift monitoring, triggered retraining, red teaming, external collaboration, cost-per-segment based threshold optimization.

11.9 Risks and Mitigation

- **Excessive customer attrition** — monitor FPR as a loss-equivalent KPI; multiply “challenge” rather than “reject”; whitelist verified behavior.

- Bias — audit of rejection/challenge rates per group; make sure sensitive proxies don't drive scores without pure risk justification.
- Attacker mapped model — do not leak detailed rejection reasons to customer channels; vary responses; monitor probing patterns.
- Late label dependency — use two-way notifications and analyst reviews as fast labels; manage explicit maturity windows.
- System failure = gate open/closed — design degradation: fallback to core rules + rescore queue; chaos exercise for payment lines.

11.10 Chapter 11 Checklist

- Map of fraud modes and policy owners per mode.
- Feature store streaming with proven point-in-time correctness.
- Inline latency budget met p99; graceful degradation tested.
- Layered ensemble (GBM-sequence-graph-anomaly rules) with cost-based threshold policy.
- Two-way real-time notifications; customer responses enter the label loop.
- Daily drift monitoring; retraining champion-challenger; red team periodically.
- Score explanations are available to analysts and auditors.

VISUAL TAKEAWAY

Fraud Control Balances Three Competing Outcomes



CHAPTER 12 — Automated Investment Advisor: Algorithm-Based Robo-Advisor

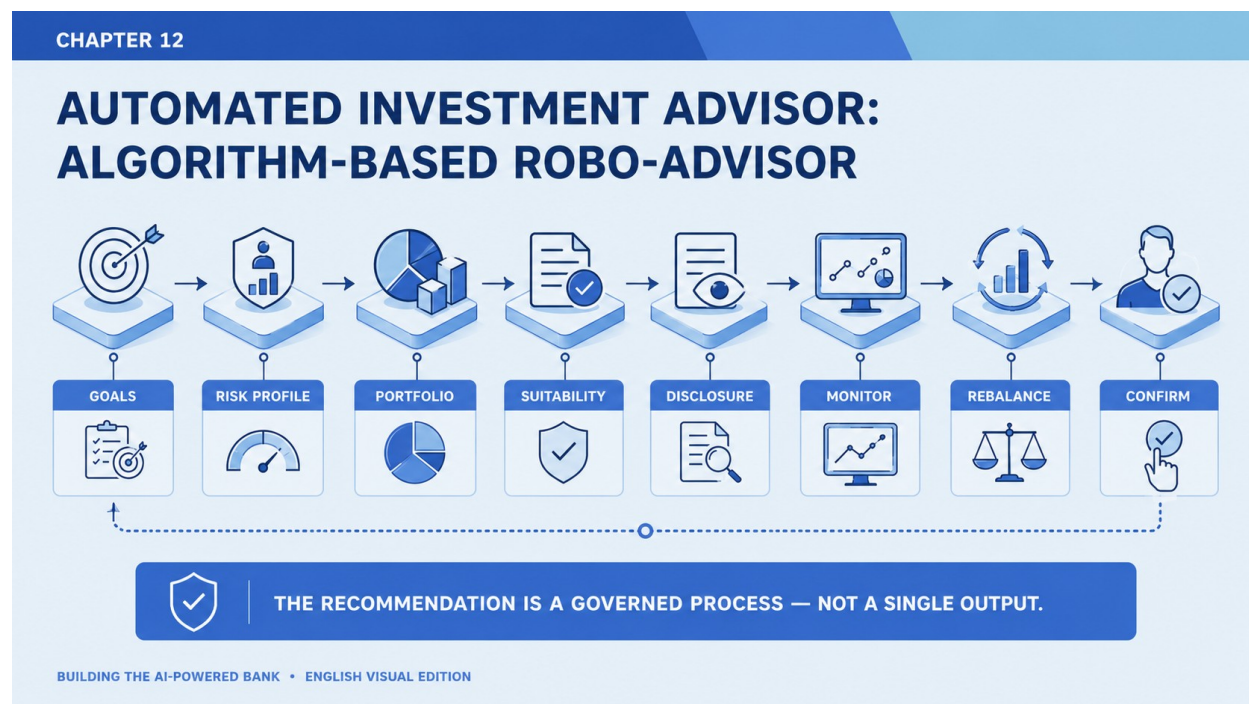


Figure 13. Automated Investment Advice Architecture — Portfolio optimization is bounded by suitability, disclosure, and ongoing monitoring.

12.1 Business Context

Quality investment advice has historically been a luxury: economical human advisors were reserved for wealthy clients, while the majority of retail clients were left to choose for themselves — often resulting in haphazard portfolios, buy-expensive-sell-cheap on emotion, or not investing at all. Robo-advisors — according to the vision map: algorithm-powered robo-advisors for selecting and managing investment portfolios — democratize that advice: with low fees and small investment minimums, each client gets a diversified portfolio that is managed disciplined: allocated according to risk profile, rebalanced automatically, and removed from emotional decisions.

For banks, the value is in the revenue quadrant (fee-based income that grows with AUM, customer retention glue, entry into long-term investment relationships) and strategic (serving mass-affluent segments that human advisors cannot reach, financial preference data that enriches 360 profiles). The key from the start: this is a highly regulated product — in Indonesia it overlaps with OJK provisions on investment advisors, APERD, and consumer protection — so compliance is not the final stage but rather the rail on which the whole system runs.

12.2 Technical Workflow: From Profile to Portfolio

VISUAL TAKEAWAY The recommendation is a governed process, not a single optimization output.

Step 1 — Profiling risks and objectives. Structured questionnaires measure three different things that are often mixed up: risk-bearing capacity (financial facts: income, liabilities, horizon), risk tolerance (psychometrics: comfort with fluctuations), and goals (20-year retirement fund vs. 3-year down payment on a house). Good practices: concrete scenario questions (“your portfolio dropped 20% in a month — what did you do?”), validation of consistency of answers, results in the form of an explainable profile, and regular updates and on-the-fly events. AI adds a layer of caution: detecting inconsistencies in answers vs actual behavior (claiming to be aggressive but panic selling at the first correction) as a trigger for profile review.

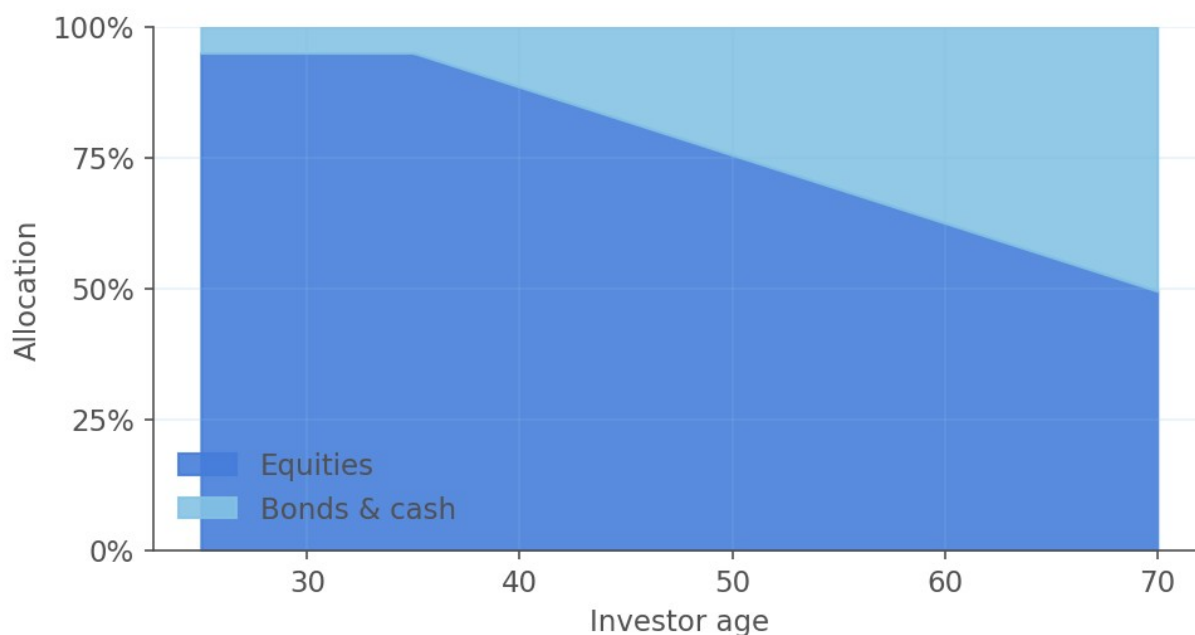
Step 2 — Portfolio construction. The theoretical foundation is Modern Portfolio Theory: maximizing expected returns at a certain level of risk through diversifying assets whose correlation is imperfect, producing an efficient frontier. Practical implementations go beyond naive mean-variance optimization — which is notoriously vulnerable to return estimation errors — with more stable techniques: Black-Litterman models (combining market equilibrium with insight), risk parity/risk contribution, bounds and regularization, resampling, and shrunk covariance estimation. The asset universe for retail is typically a low-cost index mutual fund/ETF across asset classes (domestic-global stocks, bonds, money markets, gold), mapped to a series of risk-stratified model portfolios (conservative — aggressive) that an investment committee reviews — humans set the rails, algorithms execute within them.

Step 3 — Continuous management. The engine monitors each account against its target portfolio: threshold-based rebalancing (weight deviation > x%) proven to be more efficient than pure calendaring, with cost-tax awareness (prioritizing cash inflows/outflows to balance instead of buy-sells); glide path for futures goals (risk is automatically lowered closer to the target); periodic auto-invest (rupiah cost averaging); and ongoing compliance monitoring — not just at onboarding: changes in profile, concentration, or products that are no longer suitable trigger action.

Step 4 — Modern AI layer. On top of the classic engines that are required to be solid, AI adds value to: behavioral personalization (anti-panic nudges during market corrections — timely messages based on models of customer panic tendencies, proven to save long-term returns), forecasting for optimization inputs (volatility/correlation estimates with modern time series models — safer than forecasting returns), LLM+RAG conversational assistant for portfolio education and explanation (“why do I own bonds?”) with strict fences prohibiting return promises, and experimental reinforcement learning research for allocation/execution

policies — which must be treated very carefully: RL backtests are easy to cheat (overfitting market regimes), so out-of-sample validation across regimes and strict exposure limits are an absolute requirement before touching clients' money.

12.3 Reference Architecture



An illustrative lifecycle glide path

1. Channel (app/web) → onboarding: e-KYC, profiling, digital signing, mandatory education.
2. Core services: profile & suitability engine (explicit rules + audit trail), portfolio construction service (optimization library; versioned model portfolio), rebalancing engine (event-driven on price & cash flow), order management → execution to selling agent/custodian bank/exchange according to product structure.
3. Data: market price & NAV (validated feed), customer positions & transactions, corporate actions; all on time and reconciled daily.
4. AI layer: customer behavior model, risk forecasting, LLM+RAG assistant; all via inference server and guardrails.
5. Compliance & reporting: audit trail of each recommendation/decision along with model & parameter versions, regular customer reports, regulator reporting, marketing monitoring.

12.4 Success Metrics

Growth (AUM, active customers, regular funding rate), quality of advice (tracking against model portfolio, compliance deviation = 0, total customer fees), customer behavior (retention during market corrections — good robo differentiator metrics, panic-sell rate vs. benchmarks), operations

(rebalancing accuracy, execution/reconciliation errors ≈ 0), and compliance (audit findings, complaints, marketing compliance). For investment performance: benchmark against each profile's policy benchmark — not the promise of beating the market, but the discipline of capturing market returns at low costs.

12.5 Implementation Stages

1. Licensing & product design — legal structure (MI/APERD collaboration or advisor licensing), product universe, model portfolio + investment committee, risk documents.
2. Disciplined MVP — onboarding + profiling + model portfolio + auto-invest + threshold rebalancing; without the fancy AI just yet: the right classic engine beats the wrong AI.
3. Experience layers — goal projections (Monte Carlo simulations with ranges instead of single numbers), gated LLM educational assistants, automated narrative reports.
4. Intelligence layers — personalized behavioral nudges (A/B tested), continuous compliance monitoring, risk forecasting for optimization input.
5. Scale & hybrid — hybrid tier advice (human + machine) for large clients, NBA integration (Chapter 5) and virtual advisors (Chapter 6).

12.6 Risks and Mitigation

- Mis-selling / false profiles — validated questionnaires, consistency checks, regular reviews and customer rights to understand: every recommendation can be explained in layman's language.
- Models are fragile to market regimes — avoid over-optimization on historical data; cross-crisis model portfolio stress test; The committee reviews assumptions periodically.
- Mass procyclical behavior — if all accounts rebalance simultaneously during volatility, market impact & execution queue; spread out the schedule, limit the rate, set up a high volatility mode.
- Operational failures — daily reconciliation of unit-cash positions; one-on-one control on model portfolio changes; DR for engine rebalancing.
- Marketing over-promise & LLM — prohibition on projecting returns as certainties; assistant guardrails tested; all material through compliance.
- Access bias — monitor whether product design/profiling systematically excludes certain groups from appropriate products.

12.7 Chapter 12 Checklist

- Licensing structure and roles (advisor/APERD/MI/custodian) are clearly documented.
- Risk profile: separate capacities & tolerances, validated, regularly updated.
- Verified, committee-approved, stress-tested model portfolio across regimes.
- Cost-aware threshold-based rebalancing; net daily reconciliation.

- Each recommendation is recorded with a model/parameter version and can be explained.
- LLM Assistant: RAG on official materials, prohibition on return promises, guardrails test.
- Tested extreme market behavior plans (execution rate, customer communications).

VISUAL TAKEAWAY

The Automated Advice Lifecycle



CHAPTER 13 — IVA for ATM Security and Monitoring

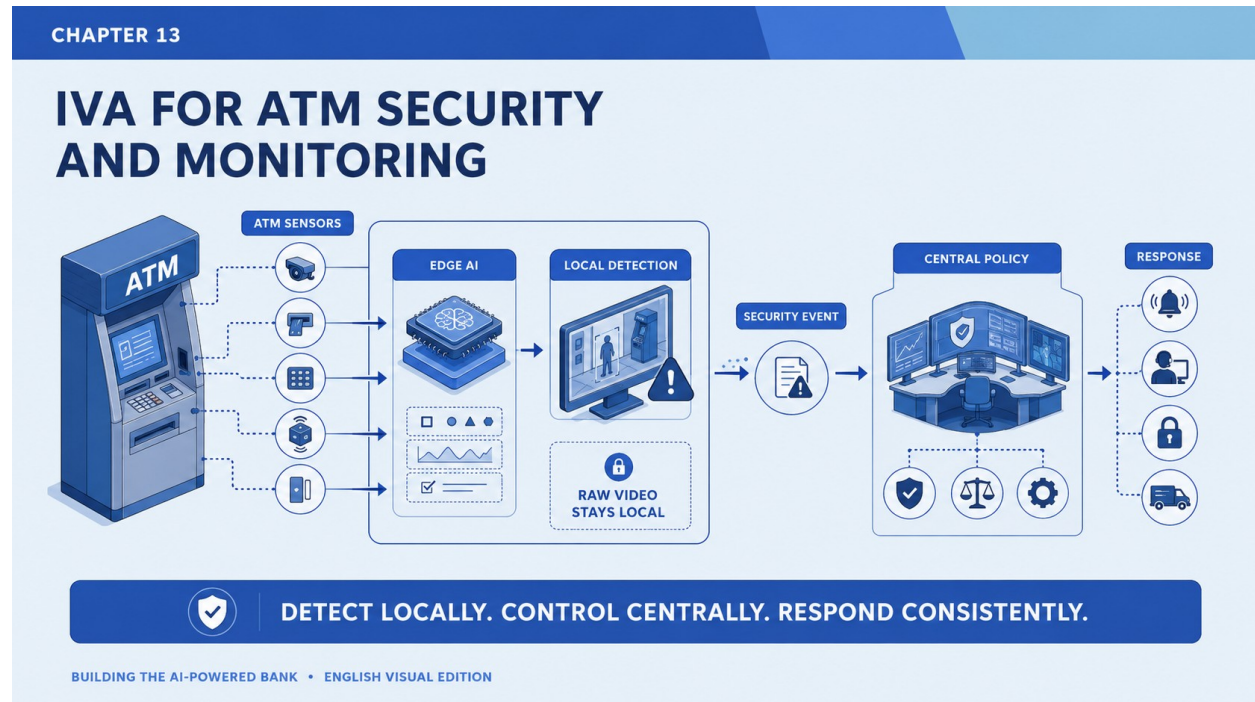


Figure 14. Embedded ATM Security Architecture — Edge intelligence detects events locally while central systems coordinate policy and response.

13.1 Business Context

ATMs are the most vulnerable point in a bank's physical footprint: operating 24 hours, often unattended, holding cash, and interacting directly with customers carrying cards and PINs. The threats are multi-layered: skimming (illegal readers in card slots + fake micro cameras/keypads to steal PINs), card trapping and cash trapping (card/money traps), physical attacks (breaking, pulling machines, gas attacks), transaction reversal fraud, vandalism, shoulder surfing and coercion against customers, to misuse of the ATM area (sleep, illegal transactions) which reduces the sense of security. According to the vision map: IVA monitors ATMs for suspicious activity, maintains customer safety during transactions, and prevents fraud.

This use case is the younger sibling of Chapter 9 with three sharp differentiators that change the design: limited environments (one-two cameras per machine, narrow and repetitive scenes — actually benefiting the model), deeply embedded computing (edges inside/next to the machine, limited heat-power-space), and multisensor fusion — the real power comes from combining visuals with transaction logs and machine sensors in real-time.

13.2 Technical Workflow

VISUAL TAKEAWAY Local detection improves latency and resilience; centralized controls preserve consistency, evidence, and accountability.

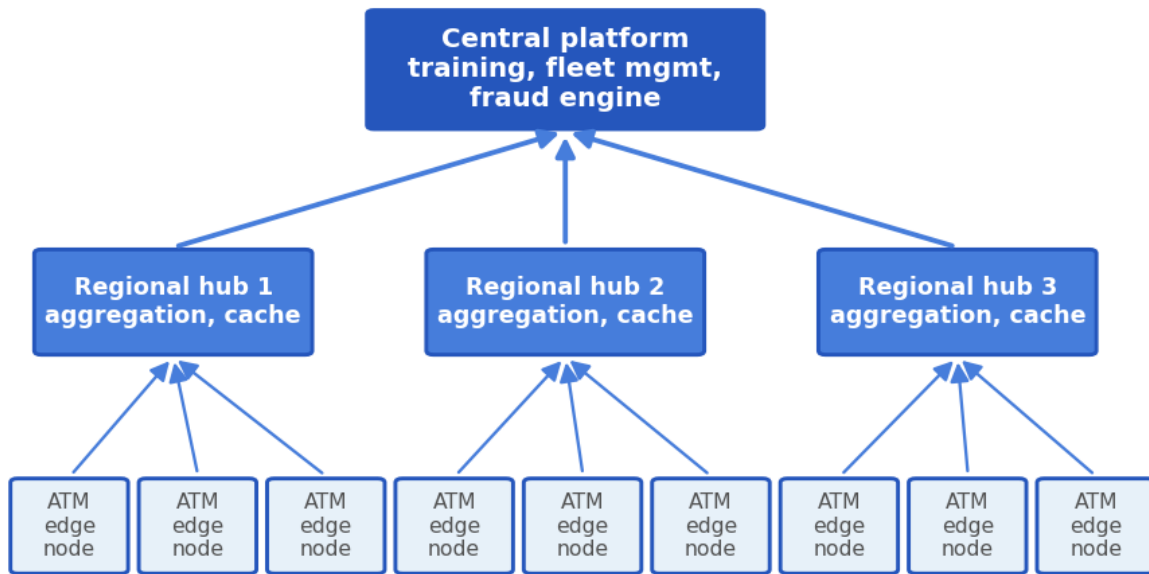
Visual anti-skimming detection. The camera faces the engine fascia; The model compares the condition of the fascia against a reference baseline per machine (detection of changes/appearance anomalies in the card slot zone, keypad and surroundings) plus detection of attached objects. The triggers are combined behavioral: people sticking to the fascia for a long time without an active transaction session (correlation with logs: no card entry) is a much stronger signal of device installation than visual alone. Typical non-visual complement: anti-skimming electromagnetic sensor in the slot — the IVA and sensor reinforce each other.

Real-time customer safety. During a transaction session (known from machine logs), the model monitors: the number of people in the transaction zone (second person close behind the customer = potential shoulder surfing/coercion → on-screen/audio warning “please keep your distance”, escalation if persistent), detection of full helmets/masks closing completely in vulnerable hours according to policy, violent gestures or customer falls (medical emergency), and abnormal repeated withdrawals under pressure (fusion of transaction patterns + visual of two people). Each policy here balances security and experience — calibrated per location and hour, with tiered response options: on-screen message → voice → central notification → officer/officer.

Asset and area monitoring. Loitering in cubicles outside of session, vandalism (hitting/damaging movements + machine accelerometer), cameras closed/distracted (automatic tamper detection), machine room doors opened outside of filling schedule (cash vendor schedule correlation), and people following filling attendants. Hourly user count metadata becomes an operational bonus: optimization of machine placement and prediction of cash requirements (forecast withdrawals per machine — a time series model that saves significant cash-in-transit costs).

Event fusion. The local correlation engine combines: visuals + transaction logs (EJ/electronic journal) + sensors (vibrate, door, anti-skim) + network status, generating severity events with complete context (“fascia anomaly + 3 consecutive failed card sessions + same person 12 minutes” ≧ sure skimming) — suppressing false alarms well below any single signal.

13.3 Embedded Edge Architecture



inference at the edge; models, policies and evidence sync with the centre

Center-region-edge topology for the ATM fleet

1. **Device:** Jetson-class embedded AI module inside an ATM cabinet (or one edge box serving multiple machines in an ATM gallery): 1-2 camera decode, INT8 quantized detection/pose model, correlation engine, encrypted record buffer.
2. **Resilience:** full operation offline — detection, local alert, record; store-and-forward event queue; watchdog and self-recovery; secure boot + encrypted disk + device identity binding (machine in public space = assume it can be touched by an enemy).
3. **Fleet center:** management of thousands of devices — zero-touch provisioning, health (camera, temperature, version), incremental OTA model/firmware updates with automatic rollback, managed fascia baseline per machine, fleet risk dashboard, command center integration and maintenance vendor ticketing system.
4. **Cyber-physical fraud integration:** ATM IVA events are published to the event bus and become a feature of the Chapter 11 fraud model — cards transacting on the machine during the suspected skimming window automatically increase their risk, enabling proactive block/monitoring before a cloned card is used. This cyber-physical fusion is this use case's biggest value multiplier.

13.4 Success Metrics

Detection (recall of skimming incidents/attacks in planned field tests, detection time, precision alerts per type), impact (skimming losses per 1,000 machines, downtime due to vandalism, customer safety

incidents), fleet (edge device availability $\geq 99\%$, OTA success, device MTTR), and efficiency (false alarms per machine per month, monitoring costs per machine vs patrol/manual methods, cash forecast accuracy when implemented).

13.5 Implementation Stages

1. Pilot 20–50 machines across location types (mall, building, drive-thru, gallery): install, calibrate baseline per machine, test with controlled attack simulation (in-house team installs dummy skimmer), measure real precision-recall.
2. Response integration — SOPs per event type (who comes, how fast), transaction log correlation, path to fraud engine.
3. Fleet industrialization — zero-touch provisioning, OTA, health monitoring; installation standards for vendors.
4. Expanding intelligence — real-time customer safety, cash predictions, placement analytics; Continuous learning from alert reviews.

13.6 Risks and Mitigation

- Public environment false alarms (shadows, bugs, new posters) — multi-signal fusion, controllable refreshed per-machine baseline, day-night adaptive threshold, periodic review per “noisy” machine.
- Edge devices compromised — full hardening (secure boot, encryption, ports down), tamper detection of the device itself, ATM network segmentation, integrity monitoring.
- Outdated model to new model — maintaining mode library (industrial intelligence), regular field testing, training-to-OTA fast track (target < 2 weeks from new mode to updated model).
- Privacy — ATM areas record customer transactions: minimization (behavioral analytics, not identity), short non-event retention, strict access, signage; strict policy of prohibiting unwarranted secondary use.
- Operational scale — thousands of devices fail slowly without first-class fleet management; invest here before adding intelligence.

13.7 Chapter 13 Checklist

- Fascia baseline per machine is managed and updated under control.
- Visual fusion + transaction log + active sensors; precision is tested with simulated attacks.
- Edge device is hardened; full offline operation; OTA with proven rollback.
- ATM events become a feature in the fraud engine (physical→cyber path) — end-to-end tested.
- SOP response per incident type with SLA; regular exercise.
- ATM area privacy policy: minimization, retention, access, signage.

PART III — ENGINEERING AND OPERATING THE AI-POWERED BANK

Nine use cases were dissected; now it's time to bring it together into one living system: end-to-end architecture, disciplined model operations (MLOps), security defenses and ethical responsibilities, regulatory compliance, and an execution strategy that takes it all from plan to reality.

PART III VISUAL MAP

Five Coordinated Planes of the Production Bank



EXPERIENCE

Channels and human workflows



INTELLIGENCE

Models, rules and agents



OPERATIONS

MLOps, SRE and FinOps



DATA & INTEGRATION

Streams, features and APIs



CONTROL

Identity, risk, privacy and audit

CHAPTER 14 — AI-Powered Bank End-to-End Architecture

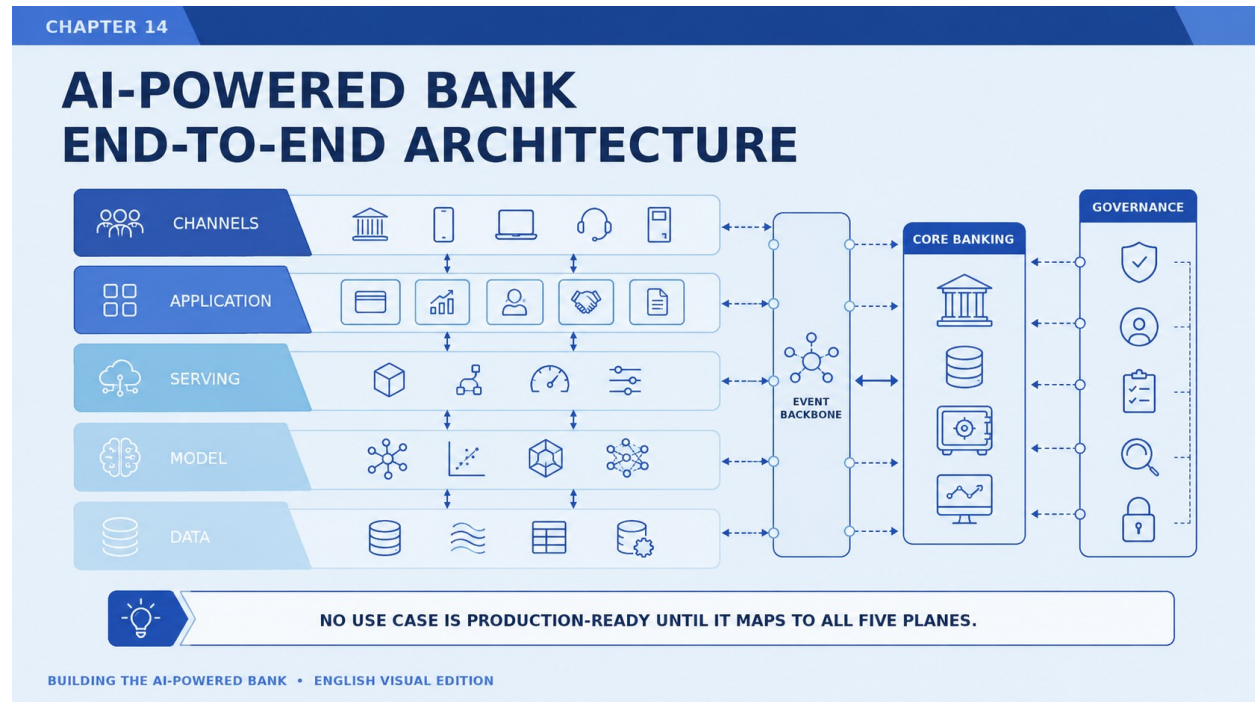


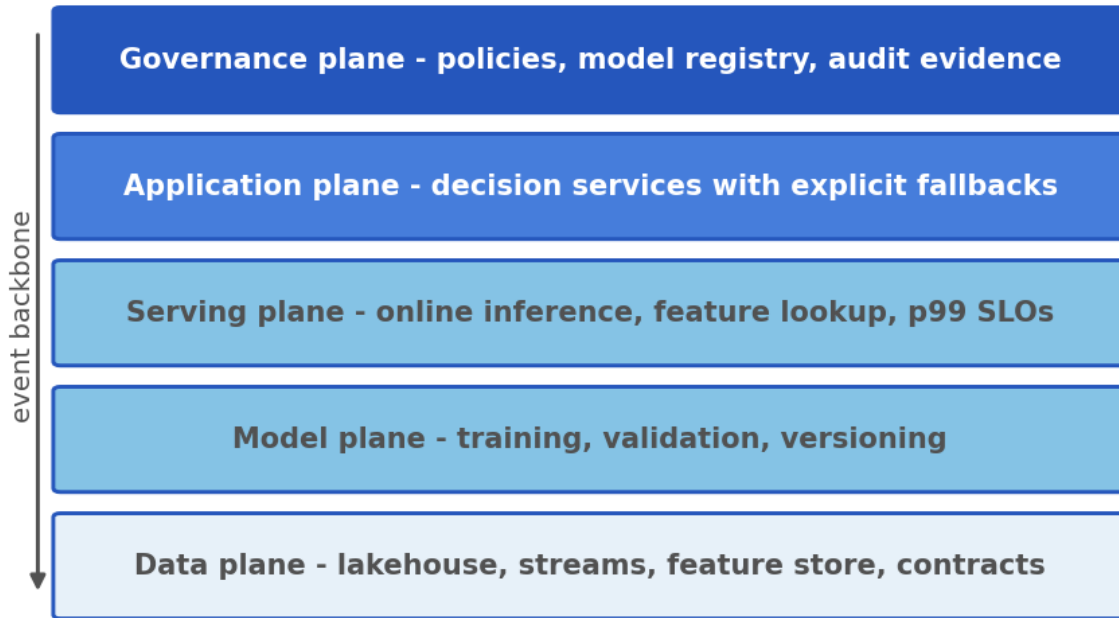
Figure 15. End-to-End AI Banking Architecture — Five coordinated planes connect channels, intelligence, data, operations, and trust.

14.1 From Nine Projects to One Platform

VISUAL TAKEAWAY Every use case should map to all five planes before it is considered production-ready.

If the nine Part II use cases were built separately, banks would end up with nine data pipelines, nine ways to deploy models, nine security standards — and integration bills that balloon every year. The right architecture reverses that ratio: 80% shared foundation, 20% use case specific logic. Notice how many requirements repeat themselves in previous chapters: unified customer identity (Chapters 5, 8, 11, 12), online feature store (5, 8, 11), event bus (all), inference server (all), edge GPU (8, 9, 13), RAG stack (6, 8, 10, 12), model registry and monitoring (all). Each of those recurring needs is a candidate for platform services.

14.2 Five Planes Blueprint



The five-plane reference architecture with its event backbone

The complete reference architecture is composed of five areas:

Data fields (details in Chapter 4): ingest batch + CDC + streaming Kafka – lakehouse medallion – semantic layer, feature store (offline/online), vector store. Data contracts between producers and consumers enforce schemas and SLAs.

Model fields: experiment workbench (managed notebook, data access via policy), GPU training cluster with scheduler, model registry (artifacts + metadata + lineage + approval status), evaluation factory (curated benchmark datasets per domain), and CI/CD/CT pipeline (Chapter 15).

Serving areas: standardized inference servers (real-time HTTP/gRPC, batch, streaming), gateway models with authentication-quota-audit, guardrails layers for LLM, caches and routers (large/small models as needed), and edge fleets (branches, ATMs) with OTA management — all exposing a uniform decision API for applications.

Application areas: Part II use cases as thin services on top of decision API — decision service NBA, dialogue orchestrator, video event engine, decisioning fraud, robo engine — plus user applications (teller, RM, command center, case analyst, customer).

Cross-cutting governance areas: IAM + attribute-based access policies, secret and key management, catalogs (data, features, models, dashboards) with connected lineage, unified observability (log-metrics-trace + model telemetry), and a policy engine (policy-as-code) that enforces rules in pipelines, not in documents.

14.3 Integration Patterns with Core Systems

Banks are not built from scratch; AI must live in peace with twenty year old core banking. Proven patterns:

- **Event backbone as a decoupler** — the core system is not randomly called by the AI; Critical changes (transactions, profile mutations) are streamed as events via CDC/adaptor to Kafka, and AI decisions come back as events/ordered API calls. The core system remains quiet; intelligence developed around him.
- **Decision service on the critical path, with degradation** — for synchronous insertion points (payment authorization calls fraud scoring), the contracts are tough: latency budgets, timeouts, and explicit fallback policies (if the AI service doesn't answer: pass-by-limit? minimum rules? queue for review?) that the business defines, is chaos-tested, and audited. AI can fail; payments should not stop without a conscious decision.
- **Strangler for gradual modernization** — legacy capabilities (fraud rules engine, legacy scoring) are not suddenly turned off; the new path runs parallel (shadow → some traffic → majority) with the continuous comparator.
- **API gateway and versioning contracts** — all cross-domain consumption via versioned APIs with backwards compatibility; no dark point-to-point integration.

14.4 Placement Topology: Center, Region, Edge

Workload placement follows data gravity and latency: central/private cloud for training, heavy analytics, large LLMs, and compliance systems; secondary region/DC for active DR and near switching latency-critical inference; edge branches for video analytics, avatar rendering, and network dropout operations; edge ATM for embedded models. Its design principle: weight raw data (video) is processed at its birthplace; only traveling events and features; models are trained centrally and distributed to the edge; and each level's ability to operate is degraded when the level above it is lost.

14.5 Resilience and Scalability

Engineering goals worth writing into architecture contracts: availability of critical decision services $\geq 99.95\%$ with graceful degradation; autoscaling inference on daily load patterns and seasonal peaks (payday, Eid, promotions); idempotency and exact-once on financial event consumption; backpressure and queues on spikes; DR tested with RTO/RPO per criticality level; and GPU capacity

is managed as a quota shared pool (Chapter 3). End-to-end latency is broken down per component budget and monitored per percentile — the p99, not the average, that customers experience.

14.6 Chapter Summary

- Build a five-plane platform (data, model, serving, application, governance); the use case becomes a thin layer on top.
- Integration into the core via backbone events and hard-contracted decision services with explicit fallbacks; modernization through the strangler pattern.
- Place the load following the gravity of the data: train centrally, serve at the edge; each level is able to stand when broken.
- Resilience is designed and tested (chaos, DR), not expected; p99 is the truth.

VISUAL TAKEAWAY

A Thin Use Case Should Sit on a Thick Shared Platform



CHANNEL & WORKFLOW



DECISION SERVICES



MODEL SERVING & MLOPS



DATA, FEATURES & STREAMS



IDENTITY, POLICY & EVIDENCE

CHAPTER 15 — MLOps: Operating Dozens of Models in Production

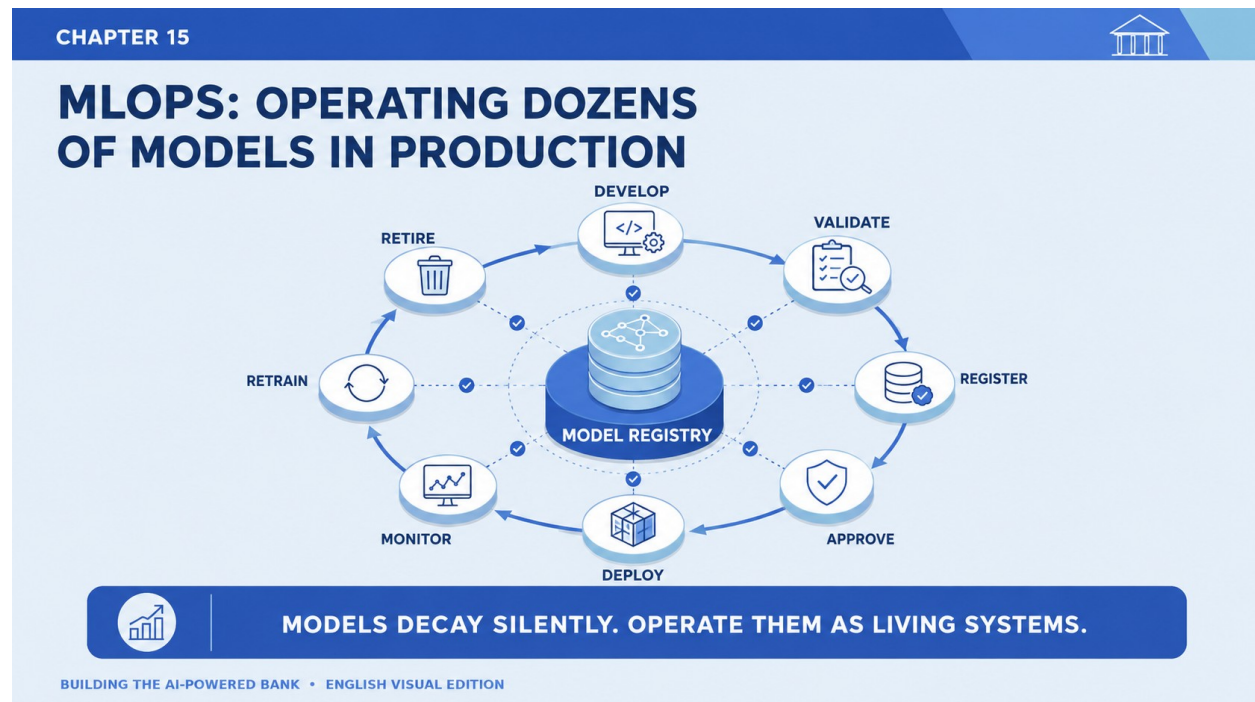


Figure 16. Managed Model Lifecycle — MLOps turns model development into a repeatable, controlled production discipline.

15.1 Why MLOps is Different from DevOps

DevOps manages one changing artifact: code. MLOps manages the interrelated triplets: code + data + model — and the model decays silently. Regular software breaks noisily (errors, crashes); models break down silently: the world shifts, input distribution changes, and accuracy slips without a single exception. MLOps therefore adds three disciplines on top of DevOps: reproducibility (every prediction can be traced to the version of code+data+model+feature that gave rise to it), model monitoring (quality of predictions as a signal of first-class operations), and continuous training (pipeline retraining as a production citizen, not an impromptu project).

15.2 Managed Model Lifecycle

VISUAL TAKEAWAY The registry is the control point linking evidence, versions, approvals, deployment, monitoring, and retirement.

Standard flow from idea to retirement:

1. Experiments — in managed workbench with automatic experiment tracking (parameters, metrics, artifacts, data versions); the best results are submitted as candidates.

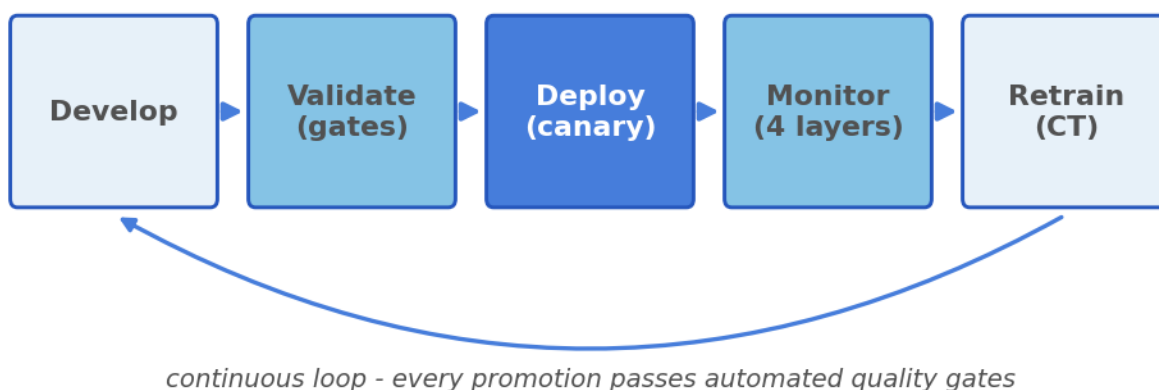
2. Automated training pipeline — not a notebook: Validated DAG (fetch data → validate data → train → evaluate → test fairness/robustness → package) that can be re-executed identically; each artifact is assigned a hash and lineage.
3. Registration and approval — registry entry model with card model (Chapter 16), evaluation results, and status; material risk models undergo independent validation (Chapter 17) before being stamped “production worthy”.
4. Progressive deployment — shadow (accepts real traffic, the decision is not used — compares silently), canary (small percentage), then full; permanent champion-challenger for dynamic domains such as fraud; one-command rollback.
5. Operations & monitoring — section 15.3.
6. Retraining & retirement — scheduled and triggered; legacy models are archived along with all their dependencies during the regulatory retention period (the ability to answer “why customer X was rejected two years ago” requires an archive of the current model+features).

15.3 Monitoring: Four Layers of Signals

1. System health — p50/p95/p99 latency, throughput, error rate, GPU saturation; usual SRE standards.
2. Data health — schema and range validation on production input; drift data: statistically measured feature distribution shifts (PSI, KS, divergence) per feature per segment; freshness of online features.
3. Model health — prediction drift (score distribution changes?), ongoing calibration, and — when labels arrive — actual performance (with label delay awareness: fraud takes weeks, credit takes months; use early proxies like analyst approval rates).
4. Business health & fairness — decision metrics (approval rate, alert rate, conversion) per protected segment and group; alarm if it shifts without explanation.

Each alarm has a runbook: who was called, a default diagnosis (corrupted data? real drift? feature bug?), and an action (rollback, emergency threshold, triggered retrain). Model incidents are treated like production incidents complete with post-mortem.

15.4 CI/CD/CT for Models



The MLOps loop: every promotion passes automated gates

- CI — in addition to code tests: feature transformation unit tests, data schema contract tests, minimum model quality tests on frozen benchmark datasets (new models must be $\geq$ threshold and not regress on critical slices), guardrails tests for LLM systems (golden questions + standard injection attacks).
- CD — standard packaging (container + model format), inter-environment promotion via registry (not manual copy), automatic progressive deployment with canary metric analysis.
- CT (continuous training) — trigger: schedule, drift over threshold, performance drop, or business event (new product, new fraud mode); The pipeline runs until the candidate is tested, then stops at the approval gate for high-risk models — full automation without humans only for policy-approved low-risk models.

15.5 Organizational Scale: Platforms and Standards

With 30–100 production models, what saves is not heroics but standards: one golden path — project templates, pipelines, and deploy patterns that the platform team supports — so that use case teams focus on the problem, not the pipeline; centralized model catalog as the sole source of status truth; clear ownership (each model has a technical owner and business owner listed); and periodic fleet reviews: which models are rotting, which are obsolete (retired — zombie models are a risk and cost), which are pending validation.

15.6 Chapter Summary

- MLOps = DevOps + data + model: full reproducibility, silent decay monitoring, retraining as a production capability.
- Managed lifecycle from experimentation to retirement, with risk-adjusted approval gateways and archives for past obligations.

- Monitor four layers: systems, data, models, business-fairness; each alarm has a runbook.
- Progressive deployments (shadow → canary → full) and champion–challenger are standard seatbelts.
- Scale is achieved through golden paths and catalogues, not through adding heroes.

VISUAL TAKEAWAY

MLOps Turns Models into Managed Services

CHAPTER 16 — Security, Model Risk, and Responsible AI

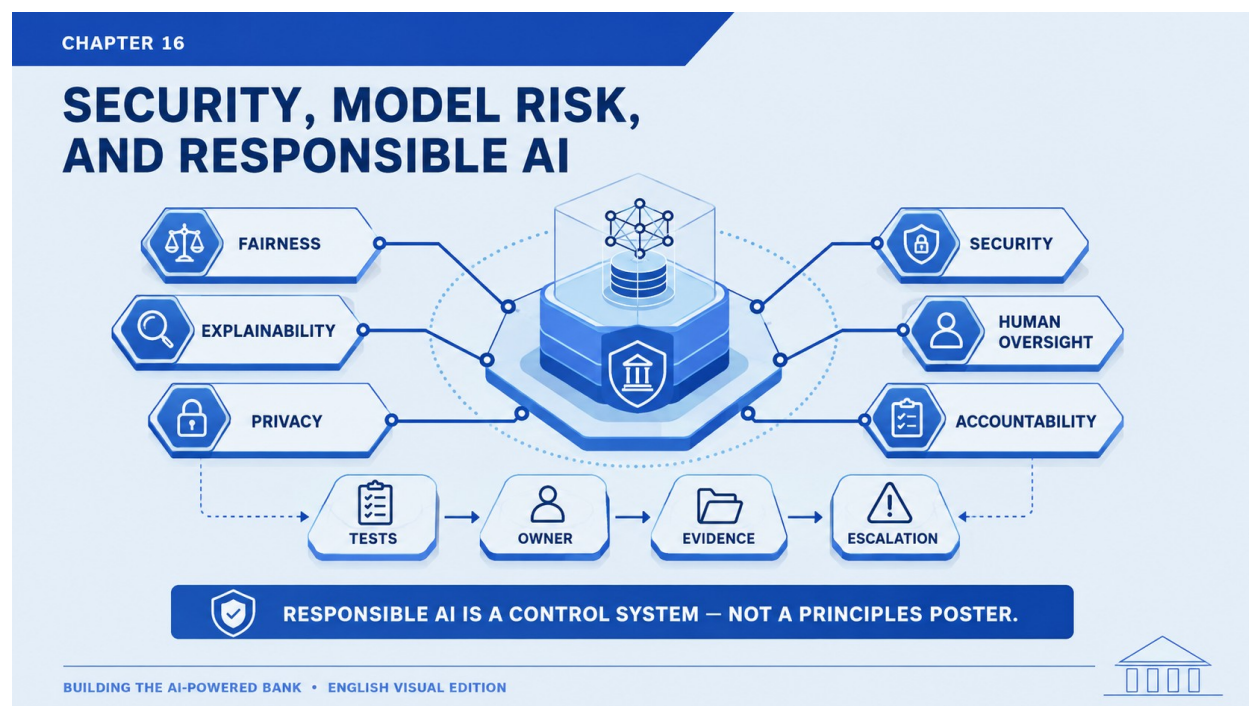


Figure 17. Responsible AI Control Pillars — Bank-grade AI requires technical and organizational controls that operate together.

16.1 AI System Attack Surface

VISUAL TAKEAWAY Responsible AI becomes operational when each principle has tests, thresholds, owners, evidence, and escalation paths.

AI systems inherit the entire attack surface of a typical IT and then add their own. AI-specific threat map that must be included in the bank's threat model:

- Evasion/adversarial examples — inputs engineered to make the model wrong: fraud transactions carved out just below the threshold; stickers/patterns that blind video detectors; text that slips past the filter.
- Data poisoning — poisoning training data or feedback loops: fraudsters “train” the system to assume patterns are normal (Chapter 11), or introduce bias through false labels.
- Model theft & inversion — stealing model behavior via bulk query (probing), or extracting training data information from the model (membership inference — a real privacy risk).
- Prompt injection & jailbreak (LLM system) — malicious instructions in user input or hidden in documents/pages captured RAG (indirect injection) that hijack behavior, leak data, or trigger unauthorized actions via function calling.
- Supply chain — contaminated pre-trained models, datasets and third-party libraries; Internal registries, artifact signing, and SBOM models are the answer.

The defenses are layered and some have been discussed contextually: input validation and sanitization, rate limiting and probing detection on model APIs, strict instruction-vs-data separation plus bi-directional filtering for LLM, minimum privileges for tools that can call models (sensitive actions are always via system authentication, not model decisions), usage anomaly monitoring, and periodic AI red teaming — internal/external teams attack fraud models, avatar guardrails, and video detectors with the latest techniques, the findings of which go into the fix backlog like regular pentests.

16.2 Model Risk Management: Bank-Grade Governance



Defense in depth around every model

Banking was lucky: the discipline of managing model risk was mature long before deep learning, crystallized in frameworks like SR 11-7 (US) and widely adopted. Its three pillars, translated into the AI era:

1. **Inventory and tiering** — all models listed; each model is assigned a risk level based on impact materiality (credit decisions and fraud = high; educational content recommendations = low); control intensity follows the level.
2. **Independent validation** — the party that did not build the model tests: conceptual feasibility (whether the approach makes sense for the problem), data quality and assumptions, performance including under slice and stress conditions, stability and sensitivity, implementation (validated model = model running in production), and documentation. For deep learning, validators add adversarial robustness tests and explanatory analysis; for LLM, evaluate guardrails and groundedness.
3. **Continuous governance** — tiered approval before production, continuous monitoring (connected MLOps Chapter 15), regular reviews, change management (material retraining = proportional revalidation), and escalation paths to the risk committee to the board. The key is stitching MRM with MLOps into one thread — not two parallel bureaucracies.

16.3 Fairness: From Principle to Test

Models trained on history will inherit historical injustices — and in banking, the stakes are credit access, fraud charges, and prices. Operational practices:

- **Define before measure** — choose fairness metrics that fit the decision context: demographic parity (equal outcome rates between groups), equalized odds (equal error rates), calibration between groups; all three cannot be met simultaneously mathematically in general — the trade-off is a policy decision that must be made consciously and documented, not left tacitly to data scientists.
- **Test includes proxies** — removing the gender/ethnicity column is not enough; features such as location, type of work, or shopping patterns can serve as proxies. Testing for disparity in results per group is mandatory even if the sensitive attribute is not used as a feature (and to test it, the attribute actually needs to be available under control in the test environment).
- **Multilevel mitigation** — pre-process (improve data representation), in-process (fairness constraints during training), post-process (thresholds per group when policy and law allow); plus product design (human review path for rejection in the gray zone).
- **Stay tuned** — justice is not a one-time certificate; drift can shift disparities (part of layer-4 monitoring Chapter 15).

16.4 Explainability

Three audiences require different explanations: the customer has the right to know the main reason for the decision that is detrimental to him in human language (“the application was rejected mainly because of the installment-to-income ratio and the last 6 months of arrears”) — not a SHAP sheet;

analysts and validators need technical feature attribution (SHAP for tabular models — consistent in theory and industry practice; attention/saliency for deep models with awareness of their limitations; counterfactuals — “what must change for the results to change” — very communicative for credit); regulators and auditors require thorough documentation: model card per model (purpose, data, performance per segment, limitations, prohibited uses) and datasheet per dataset. The architectural principle: explanations are calculated and stored with decisions at the time they are made — reconstructing them later from a model that has changed versions is an audit nightmare.

For high-impact decisions, model choice is also part of explainability: if a glass-box model (scorecard, monotonic-constrained GBM, modern GAM) achieves black-box equivalent performance, choose explainability — small accuracy differences are rarely worth the loss of auditability.

16.5 Meaningful Human Oversight

“Human in the loop” is easy to say and easy to fake: humans approving 500 machine recommendations per day with a 99.8% acceptance rate is not oversight, but a rubber stamp. Meaningful monitoring design: sort decisions based on risk (full automation for low-monitored risk only; gray zone and high impact human mandatory); give supervisors enough information and enough time (explanation, context, realistic workload); measure supervision quality (correction rate, inter-inspector concordance, blind sample audits where supervisors assess without looking at machine scores — detecting automation bias); and protect the authority of resistance: the overseer who never differs from the machine, or who is punished when he differs, is not the overseer.

16.6 Operational Ethics Framework

The bank's AI principles (fair, transparent, secure, accountable, private) only live when given teeth: a cross-functional AI ethics/risk committee with termination authority; mandatory impact assessment before new use cases (purpose, affected parties, rights risks, mitigation — integrated with DPIA); explicit prohibition list (common examples: social scoring, inference of sensitive attributes without basis, manipulation of psychological vulnerabilities for sales, mass biometric identification without legal basis); complaint and correction channels for customers regarding automatic decisions; and annual training for AI system developers and users. An ethic that stops nothing, never, is suspect of mere decoration.

16.7 Chapter Summary

- AI systems add new attack surfaces — evasion, poisoning, theft, prompt injection, supply chain; the answer is layered defense plus routine AI red teaming.
- The legacy of banking MRM (inventory-tiering, independent validation, sustainable governance) is the golden framework — stitched into one thread with MLOps.

- Fairness demands trade-off-aware metric choices, proxy testing, multilevel mitigation, and continuous monitoring.
- Explanations are tailored to the audience and stored with the decision; for high impact, consider a glass-box model.
- Human supervision must be measured in quality; ethics must have the authority to stop.

VISUAL TAKEAWAY

Responsible AI Is a Control System, Not a Principle Poster

CHAPTER 17 — Regulation and Compliance: Operating in a Regulated Environment

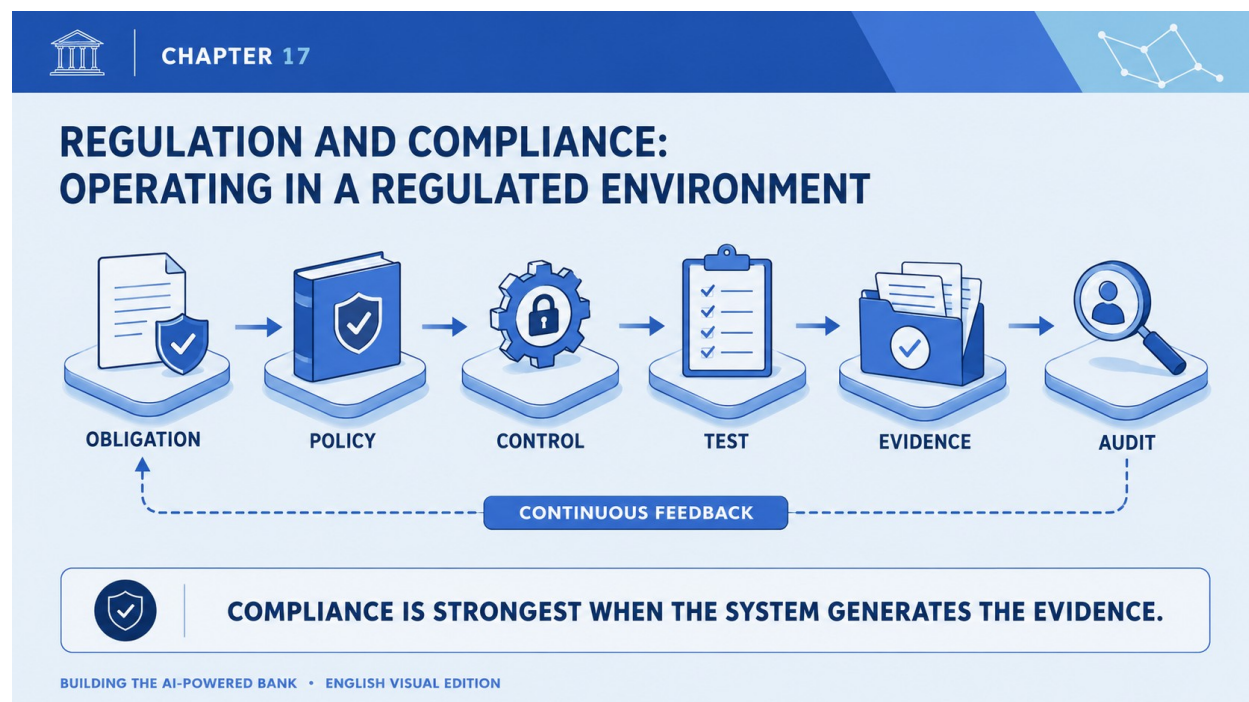


Figure 18. From Regulation to Technical Evidence — Translate obligations into controls that can be tested and demonstrated.

17.1 Regulatory Map for Bank AI in Indonesia

VISUAL TAKEAWAY Compliance is strongest when evidence is generated continuously by the system—not assembled at the end.

Every use case in this book touches on at least one regulatory regime. A summary map for the Indonesian context (details of the provisions continue to evolve — treat this section as a regional map, and always consult the current regulatory text with your legal advisor):

- **Personal data protection** — Law no. 27 of 2022 (PDP Law) determines the basis of processing (consent and other legal grounds), data subject rights (access, correction, deletion, withdrawal of consent, objection to automated decisions), controller obligations (purpose limitation, minimization, security, leak notification), as well as data transfer provisions. Immediate implications for AI: clarity of the legal basis for each processing (personalization ≠ operational), special attention to automated decisions with significant impact (credit, denial of service) that require explanations and objection paths, and governance of biometric and other specific data.
- **OJK regulations** — umbrella for the implementation of information technology by banks (including IT risk management, outsourcing and cloud computing), general risk management, governance, financial services sector consumer protection, digital banking provisions, as well as capital markets sector

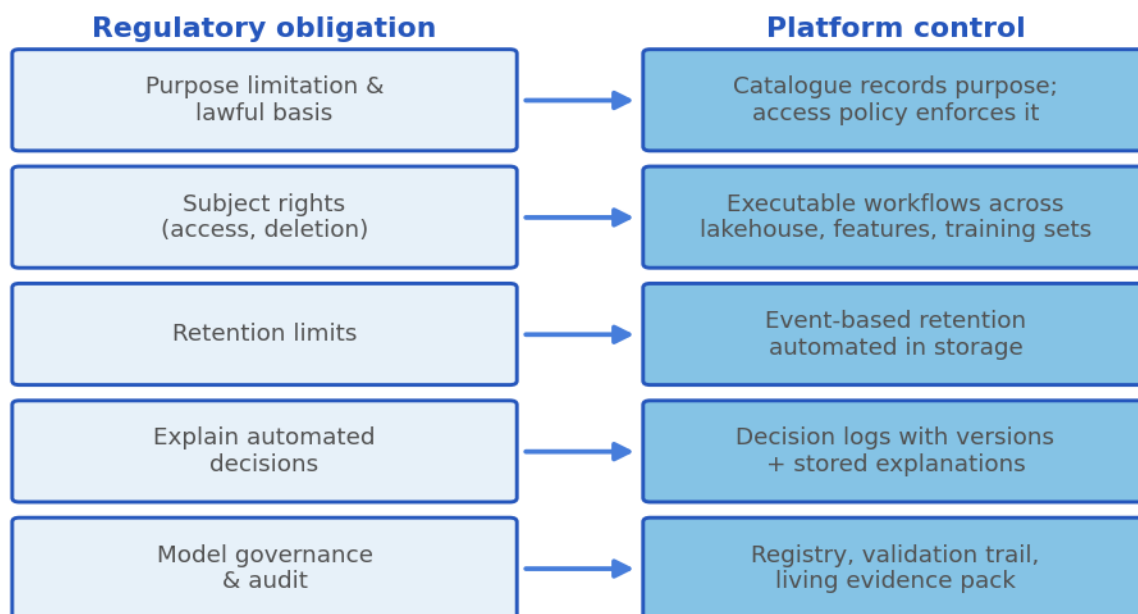
provisions for investment advice/management activities (directly relevant to Chapter 12). OJK has also issued policy guidelines/directions regarding the use of AI in the financial sector which emphasizes governance, risk management and human responsibility.

- Bank Indonesia regulations — payment systems (relevant to real-time fraud pathway Chapter 11), transaction security standards, and payment infrastructure implementation.
- APU-PPT (AML/CFT) — anti-money laundering and terrorism financing prevention program obligations: customer due diligence, transaction monitoring, reporting suspicious transactions to PPATK — an area where AI is actually the main compliance tool (transaction monitoring, sanctions screening, mule network detection) as well as an object of supervision (the monitoring model must be accountable).
- Other provisions — bank confidentiality, employment (impact of automation), and industry standards such as PCI DSS for card data.

17.2 The Global Landscape that Shapes Practice

Indonesian banks that operate or partner across borders — and anyone who wants to follow best practices — need to understand global currents: GDPR (European Union) as a data protection mecca that popularized the right to automated decision explanations; The EU AI Act as the first comprehensive AI regulation with a risk-based approach — AI systems for creditworthiness are classified as high risk with obligations for data governance, technical documentation, transparency, human oversight, resilience and conformity assessment; AI risk management frameworks such as the NIST AI RMF and AI management system standards (ISO/IEC 42001) which serve as audit references; banking model supervision principles (SR 11-7 tradition, guidance of European and Asian authorities); as well as Basel standards for risk models used for capital calculations. The broad pattern of all these regimes is uniform and is already the backbone of Chapter 16: risk-based, demanding documentation and explanation, requiring human oversight for high impacts, and requiring system-wide monitoring.

17.3 Translating Regulations into Technical Controls



From obligation to control: compliance engineered into the platform

Compliance that persists is compliance that is embedded in the platform (compliance by design), not one that is pursued ahead of an audit. The core mapping table:

Regulatory Demands	Technical Controls on the Platform
Basis & purpose of processing	Data catalog + purpose tags + purpose-based access policies
Minimization & retention	Tokenization since ingest; automatic retention schedule
Data subject rights	Cross-system access/deletion APIs + execution evidence
Explanation of automated decisions	Explanations retained with the decision (Chapter 16)
Appeal and human review	Case workflow with SLA and correction authority
Model auditability	Registry + lineage + versioned model/feature archive

Regulatory Demands	Technical Controls on the Platform
Security	Encryption, IAM/ABAC, comprehensive access logging
Incident and breach reporting	Detection + time-bound notification workflow
Outsourcing / cloud	Processor register, contract clauses, location and key controls
Continuous oversight	Chapter 15 monitoring and periodic committee reporting

17.4 Dealing with Supervisors and Audit

Practices that differentiate “audit-safe” banks from “watchdog-trusted” ones: involve compliance and risk management functions from the time of use case design (a permanent representative on each use case team — not the final checkpoint); maintain a package of live evidence per material risk model (model cards, validation results, monitoring logs, committee decisions) that can be submitted at any time without panic; run independent tests against the reference framework (NIST AI RMF / ISO 42001) before outside parties do; and build proactive relationships — regulators in many jurisdictions provide sandboxes and consultation forums for innovation; coming in first with designs and mitigations is much better than being called in later with findings.

17.5 Chapter Summary

- Indonesia's regulatory landscape for banking AI centers on the PDP Law, the OJK framework (IT, risk management, consumer protection, capital markets), BI provisions for payments, and the APU-CFT regime.
- Global currents (GDPR, EU AI Act, NIST AI RMF, ISO 42001, MRM tradition) are standardizing the direction: risk-based, documented, explained, human-supervised, lifetime-monitored.
- Translate demands to platform controls (compliance by design); maintain proof-of-life packages; and make the regulator a dialogue partner from the start.

CHAPTER 18 — Implementation Strategy: From Roadmap to ROI

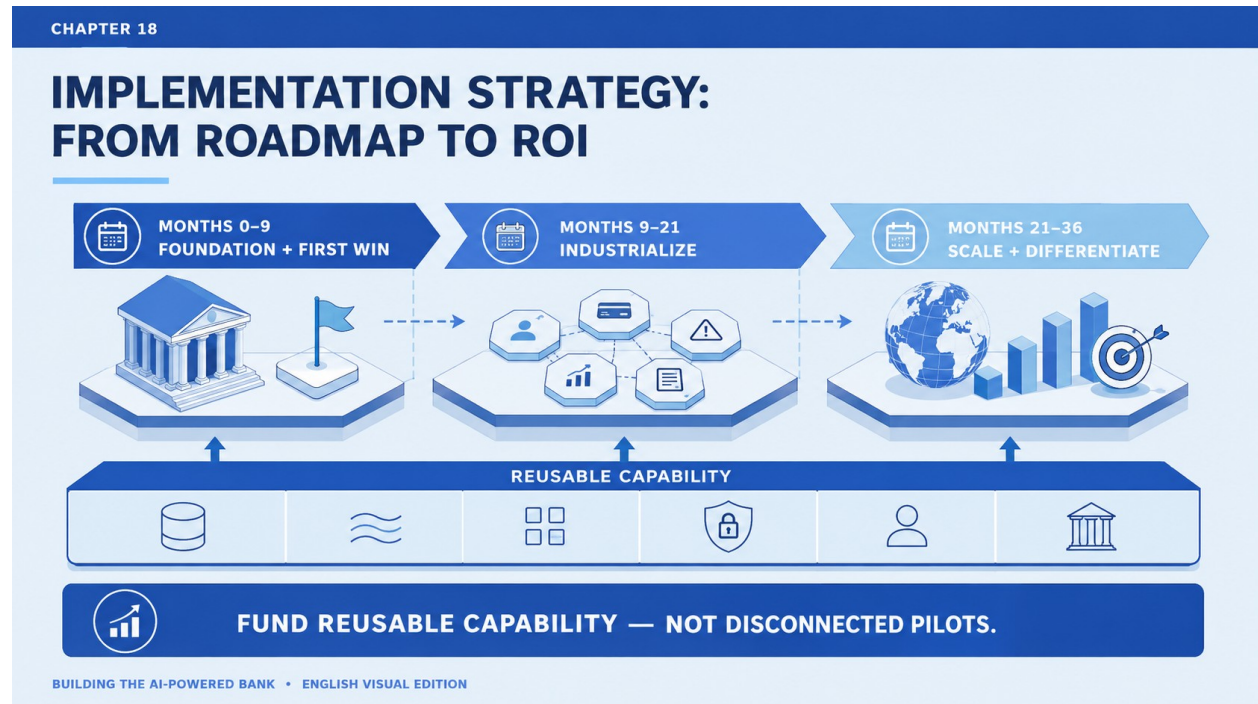
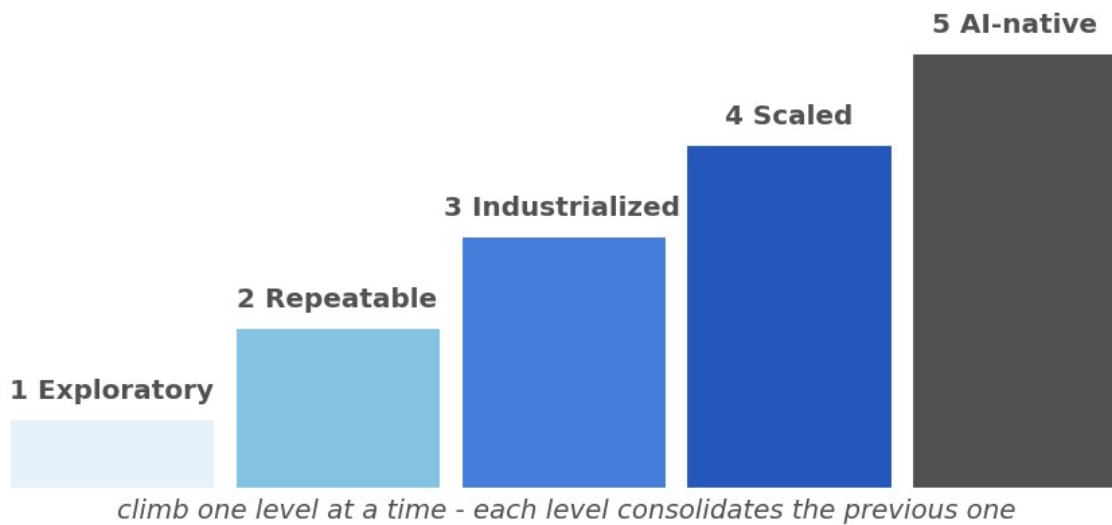


Figure 19. 36-Month AI Transformation Roadmap — Sequence foundations and controls before scaling the portfolio.

18.1 Assessing Maturity: Know Your Position Before You Move



Five levels of AI maturity - climb one at a time

A good strategy starts from position honesty. A practical five-level maturity model:

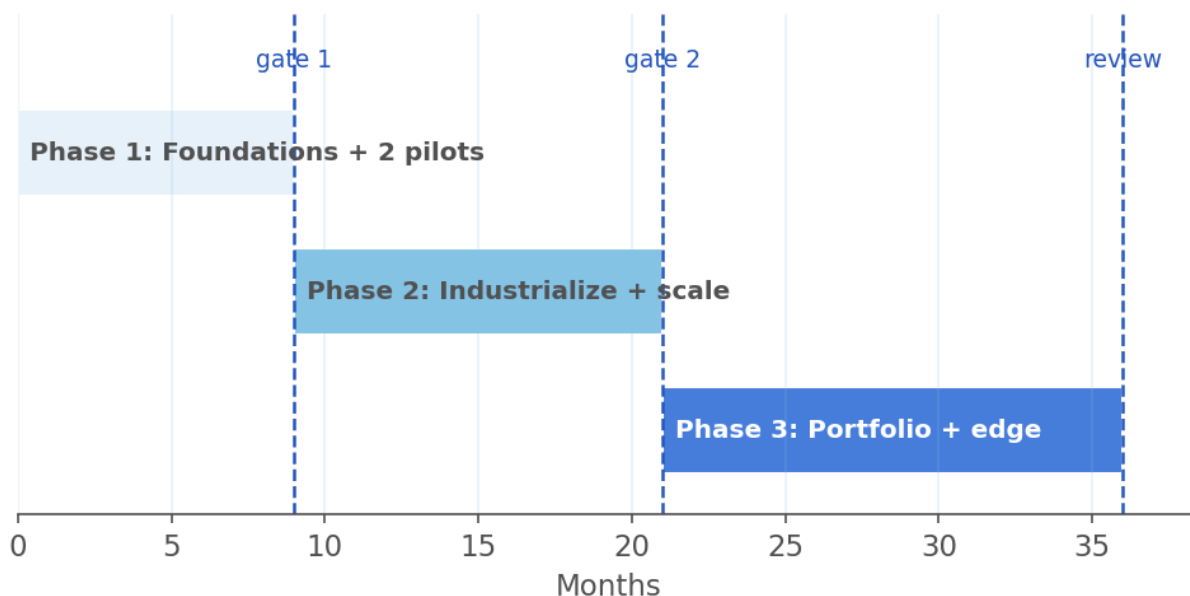
1. **Exploratory** — sporadic PoCs, data silos, no standards; Success depends on the individual.
2. **Active** — some isolated production models; start a data team; there is no shared platform yet.
3. **Operational** — initial AI-data platform, basic MLOps, 5–15 monitored production models; model governance is starting to become formal.
4. **Systemic** — mature platform used across domains, feature store and registry are standard, integrated MRM, dozens of models; AI enters business planning.
5. **Transformative** — AI-based decisions become the default way of working; quick-cheap experiments; Use case portfolios are managed like investment portfolios.

Assess yourself with evidence (number of production models, idea-to-production time, monitoring coverage, audit findings), then design a one-level jump — jumping two levels at once almost always fails because the organization doesn't have time to grow its muscles.

18.2 Selecting the Order of Use Cases

Prioritize along two axes — value (financial + strategic impact) and feasibility (data readiness, integration complexity, regulatory burden, organizational readiness) — plus one architectural consideration: the foundation on which it is built. The usual sequence of success in banks: starting from fraud/AML enhancement and risk analytics (Chapters 10–11) because the data is the most ready, the measurable value is clear, and forces the birth of a streaming + feature store foundation; further personalization (Chapter 5) and text-voice assistants (Chapter 8) that lay the groundwork; then physical computer vision (Chapters 9, 13) which requires edge investment; and finally immersive experiences and twins (Chapters 6, 7) which are the most exploratory — unless there are specific business drivers that change the order. Manage it as a three-horizon portfolio: majority of the budget to proven use cases, some to expansion, a few to explorative bets — and dare to kill off the ones that don't prove themselves.

18.3 36-Month Reference Roadmap



Three gated phases over thirty-six months

VISUAL TAKEAWAY The roadmap should fund reusable capability, not a collection of disconnected pilots.

Phase 1 (months 0–9): Foundation + first wins. Establish a core team and governance; build thin slices of the platform (core lakehouse, streaming for one domain, feature store v1, standard deploy path); deliver one end-to-end use case to production (instead of three PoCs) with measurable business metrics; set basic MRM and privacy policies. Key outputs: proof of value + embryonic platform + organizational trust.

Phase 2 (months 9–21): Industrialization. Expand to 3–5 use cases on the same platform; mature MLOps (four layers of monitoring, retraining, champion–challenger); formalize independent validation; build edge capabilities if video/ATM joins the wave; start an organization wide AI literacy program. Key output: marginal cost of new use cases drops drastically; The first audit is passed with the evidence organized.

Phase 3 (months 21–36): Scale and differentiation. Portfolio of 10+ use cases; cross-channel orchestration (NBA everywhere, fraud-ATM cyber-physical fusion); exploratory use cases (immersive, twin) tested with discipline; FinOps and fleet review models become a rhythm; long-term talent strategy in place. Key output: AI as an institutional capability, not a program.

18.4 The First 90 Days: From Mandate to Production

The first ninety days should not promise enterprise transformation. They should prove that one material decision can move safely from data to production while leaving behind reusable capability.

The clock starts only after an executive sponsor, product owner, model owner, and independent risk counterpart accept named accountability.

Window	Purpose	Required Evidence
Days 0–15	Outcome and authority	Define the customer or risk outcome, baseline, economic owner, decision rights, risk tier, and a single production definition of done.
Days 16–30	Data and control feasibility	Test point-in-time data, consent and lawful use, lineage, access, integration paths, latency, validation scope, and operational ownership.
Days 31–60	Build and shadow	Create the thin end-to-end slice; establish an evaluation set; run shadow traffic; instrument business, model, risk, and operational metrics.
Days 61–75	Controlled production	Release to a narrow cohort with human approval, transaction limits, rollback, incident response, and daily review of errors and overrides.
Days 76–90	Production-readiness decision	Approve, redesign, or stop using evidence. Transfer ownership to operations and record the platform components that the next use case can reuse.

Stage Gates: Evidence Before Enthusiasm

Gate	Decision	Minimum Evidence	Stop Condition
G0	Outcome	Named owner, baseline, target, control group or credible counterfactual	No measurable decision or accountable owner
G1	Data	Lineage, quality profile, access approval, consent/legal basis, leakage test	Critical data cannot be used lawfully or reproducibly
G2	Model	Out-of-time validation, calibration, segment tests, explainability, limitations	Performance is unstable or harm cannot be bounded
G3	System	Load, resilience, security, abuse tests, observability, fallback and rollback	No safe degradation path or incomplete audit trail
G4	Operations	Runbook, service owner, on-call, incident severity model, change controls	Nobody accepts production ownership
G5	Value	Incremental outcome, adoption, risk impact, total cost, residual risk	Value is not incremental or residual risk exceeds appetite

18.5 The First 12 Months: From One Model to a Portfolio

A credible first year compounds capability. Each quarter must produce both a business outcome and a reusable platform or governance asset. If the second use case costs as much and takes as long as the first, the bank has delivered projects rather than an institutional capability.

Quarter	Business Milestone	Capability Left Behind	Board Evidence
Q1 — Prove	One narrow production use case	Baseline, stage-gate pack, monitored thin slice, accountable operations	One outcome is demonstrably safer or better
Q2 — Industrialize	A second use case reuses the stack	Standard deployment path, registry, feature reuse, validation templates	Idea-to-production time begins to fall
Q3 — Scale	Three to five governed capabilities	Shared monitoring, FinOps, platform SLOs, domain squads, control automation	Marginal delivery cost and incident rate remain controlled

Quarter	Business Milestone	Capability Left Behind	Board Evidence
Q4 — Govern the portfolio	Fund, retire, and sequence models as assets	Quarterly fleet review, benefit realization, technical-debt register, 24-month roadmap	Capital follows evidence; weak use cases are stopped

The Minimum Accountability Map

Role	Accountability	Non-Delegable Evidence
Executive sponsor	Accountable for outcome and risk acceptance	Funding, priority, conflicts, stop decision
Product owner	Accountable for workflow adoption and benefit realization	Baseline, process design, user feedback, KPI
Model owner	Responsible for model behavior through its life cycle	Data, evaluation, limitations, monitoring, change
Platform / service owner	Responsible for reliability and recoverability	SLOs, capacity, deployment, observability, incident response
Risk / compliance / privacy	Independent challenge and control interpretation	Risk tier, validation scope, legal basis, approvals, evidence
Security	Responsible for threat model and technical assurance	Identity, secrets, abuse testing, logging, containment
Internal audit	Provides independent assurance, not project sign-off	Tests governance design and operating effectiveness

18.6 Organization and Talent

Proven organizational models for bank scale: hub-and-spoke — the center (hub) owns the platform, standards, MRM and rare specialist talent; The use case squad (spoke) contains a mix of centers + domains (product business owners, data scientists, ML engineers, data engineers, designers, risk/compliance reps) who are end-to-end responsible for their business metrics. The roles that are most often underinvested in are not data scientists, but rather: machine learning engineers / platform engineers (who take models to production), data engineers (who make things possible), AI product managers (who keep problems real), and independent model validators. Realistic talent strategy in a tight market: build-from-within (in-house engineer conversion program — bank IT engineers who know their systems are the best raw material), selective hire for leverage roles, university partnerships, and retention through compelling problems plus managerial-equivalent technical career paths.

18.7 Build, Buy, or Partner

The decision framework is per layer: buy/subscribe to non-differentiating commodities (underlying infrastructure, generic MLOps tools, generic foundation models); build a differentiator and get to the heart of customer-data (fraud/credit/NBA core features and models, knowledge base, core integration); partner for domain-specific acceleration (IVA vendors, digital human platforms, data providers) — with contractual terms that preserve sovereignty: ownership of data and derived

models, audit rights, portability (output in open formats), SLAs and local support, and a ban on using bank data to train other parties' models without explicit permission. A classic trap to avoid: buying a black box “AI solution” for a high-stakes decision that can't be validated — regulators (and Chapter 16) don't take “our vendors know” for an answer.

18.8 Program Economics: Honest TCO and ROI

Cost side (3–5 year TCO): infrastructure (GPU/cloud, storage, network, edge), software and subscriptions, people (largest and most often underestimated — including platform and validation teams), data (acquisition, labeling, quality), integration and process changes, compliance and audit, and ongoing operations (the model is never “done”). Benefit side: map to Chapter 1's four quadrants with measurement discipline — benefits are claimed only when proven to be incremental (control group, A/B, or trend-corrected before-after baseline), with the business owner signing off on the numbers, and tracked post-implementation (not just in the proposal). Healthy practices: business case per use case + quarterly benefits realization report to the steering committee; use cases that are not even close to promises in two quarters are restructured or discontinued. An allocation rule of thumb that often surprises newcomers: for every dollar a model costs, set aside a few dollars for data, integration, and process changes — that's where the real value is created or lost.

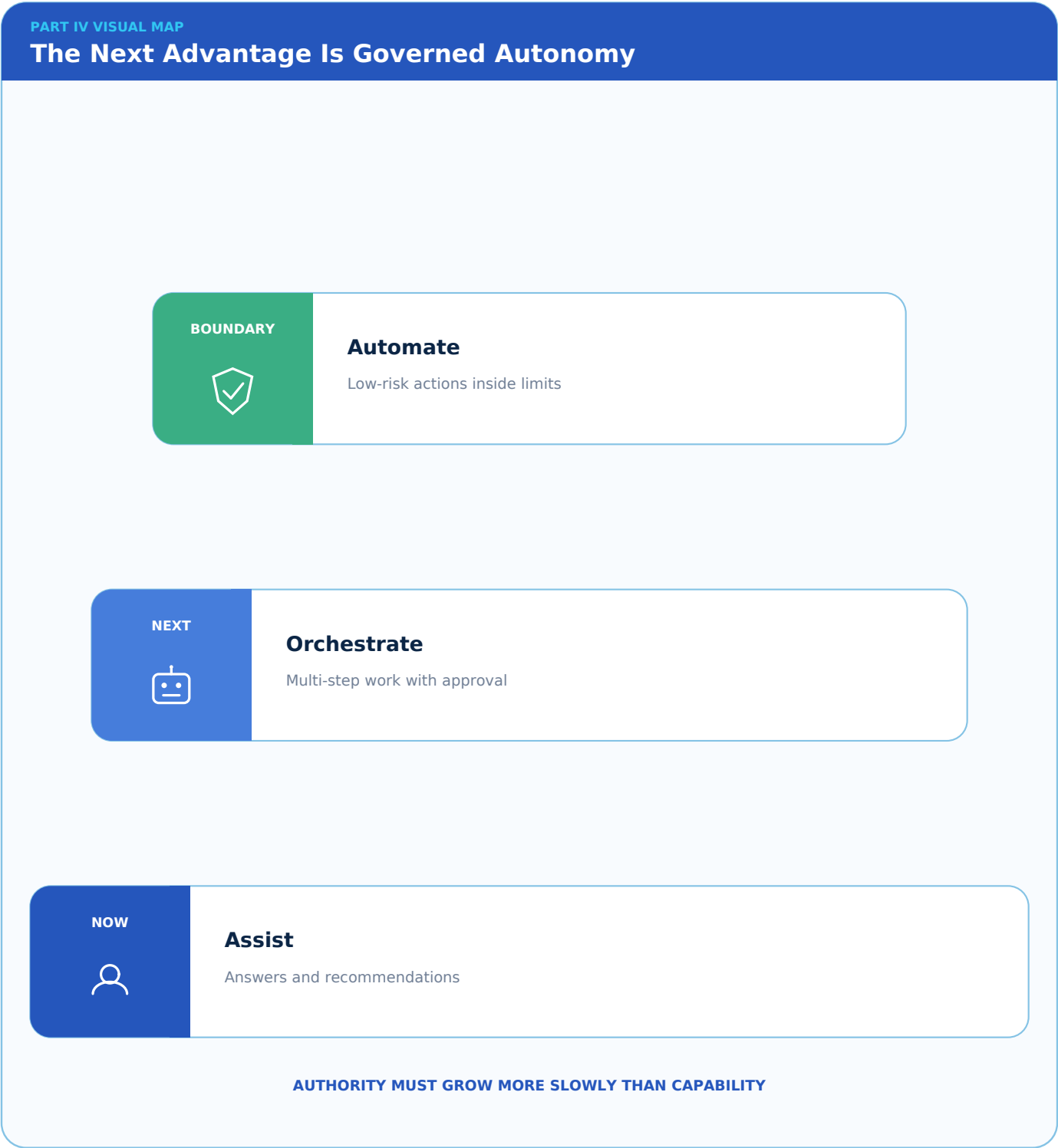
18.9 Change Management: The Human Side

AI programs rarely fail because the algorithm is unavailable; they fail when the workflow, incentives, or trust model is wrong. Involve tellers, analysts, operators, and control functions in design because they know where the real process breaks. Begin with advice before automation, explain what the system can and cannot do, preserve an effective way to challenge decisions, and fund genuine reskilling when roles change. Measure adoption, override quality, and time saved as seriously as model accuracy. Change management is product design from the first day, not communication at the end.

18.10 Chapter Summary

- Assess maturity honestly and advance one evidence-backed level at a time.
- Prioritize by value × feasibility × reusable foundation; stop weak use cases quickly.
- In year one, prove one production thin slice, industrialize reused controls, then govern the portfolio.
- Use hub-and-spoke ownership and invest in engineering, validation, operations, and change—not only data science.
- Buy commodities, build differentiation, demand auditability, and measure only incremental value.

PART IV — LESSONS AND THE FUTURE



CHAPTER 19 — Case Studies and Patterns from the Field

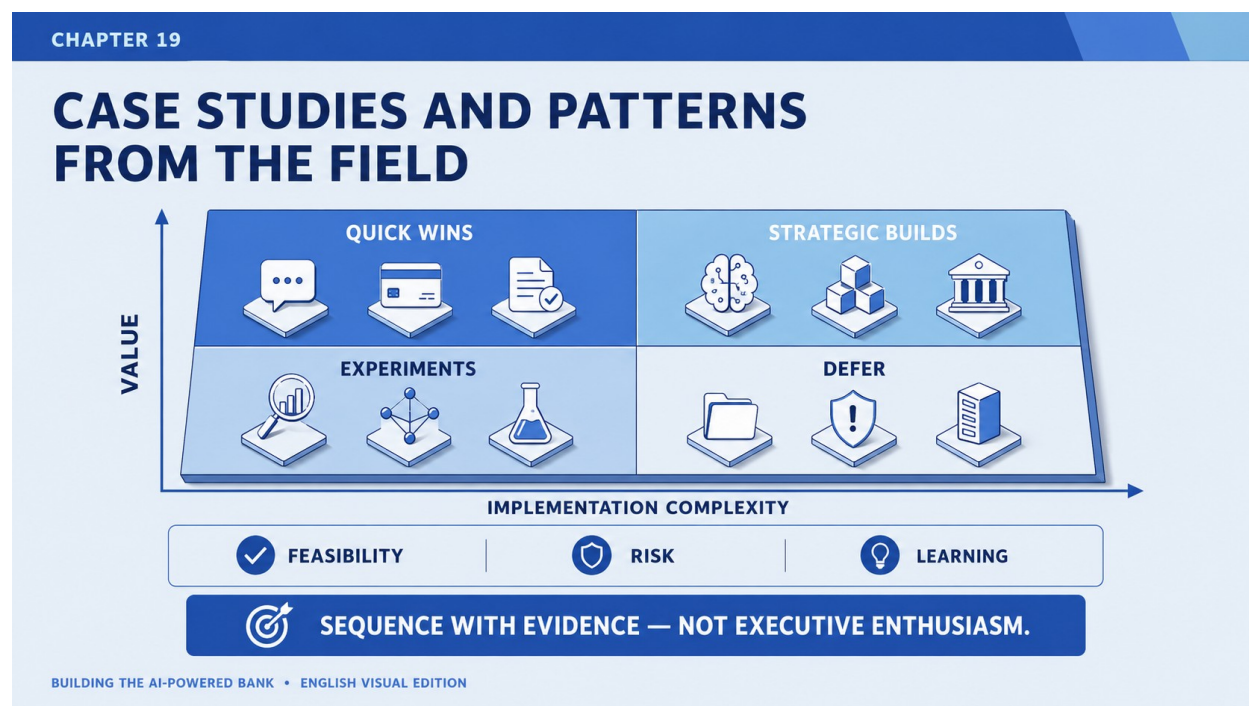


Figure 20. Value vs. Implementation Complexity — Illustrative portfolio chart for sequencing—not a universal ranking.

This chapter presents three composite case studies — abstracted from real, recurring patterns in the industry, with names and numbers disguised/illustrated — and then summarizes the most frequently occurring failure anti-patterns and success patterns. The goal is not to imitate, but rather to recognize dynamics that you will almost certainly encounter.

19.1 Case Study A: Large Universal Bank — Fraud as a Pathfinder

VISUAL TAKEAWAY Sequence work using value, feasibility, risk, and learning potential—not executive enthusiasm alone.

Situation. Banks with tens of millions of customers face a surge in digital channel fraud; Legacy regulatory systems produce an abundance of alerts with very low precision, sink analyst teams, and increase customer attrition.

Action. Instead of a “fraud project,” management agreed to build the foundation: streaming Kafka from switching, a feature store with real-time velocity features, and a decision service with tens of milliseconds latency — then gradient boosting models replaced the majority of the rules, followed by behavior sequence models and entity graphs for mule. Two-way notifications to customers are installed, and every analyst decision becomes a label.

Results (illustrative, patterns commonly reported by industry). In four quarters: fraud losses fell by double digits, precision alerts rose multiple times to normalize analyst backlogs, and false declines fell — satisfaction rose while security tightened. Even more strategic: the feature store and real-time pipeline that emerged from this project became the foundation of the NBA and digital assistants of the following years — the cost of the second and third use cases was a fraction of the first.

Lesson. The right first use case is the one whose value is undeniable and forces the birth of the foundation. Loop labels have been from day one our most valuable silent asset.

19.2 Case Study B: Mid-Sized Bank — An Avatar that Nearly Failed

Situation. Banks launch virtual assistants with great enthusiasm; the internal demo is stunning. Three months post-release: low containment, complaints of "inconsequential answers" going viral on social media, and the team facing the pressure of shutting down the service.

Diagnosis. Classic: knowledge base built from outdated documents without an update pipeline; there is no automatic evaluation (golden questions) so regression is not detected; thin guardrails; and escalation to humans removes context so the customer repeats the story — exactly the most annoying moment.

Recovery. The team stopped adding features and rebuilt disciplines: re-curated the knowledge base with per-topic content owners and update SLAs; RAG with obligatory citations and honest "I don't know" answers; hundreds of gold questions in CI; red team injection; contextual handover; as well as daily sample review by the quality team. The scope is deliberately narrowed to topics that are mastered, then expanded gradually based on failure data.

Results. Six months later containment was healthy and satisfaction scores surpassed the phone channel for topics covered. **Lesson:** in the LLM system, excellence is not in the model that everyone uses, but in the discipline of evaluation, knowledge curation, and graceful failure design. A narrow and honest second launch beats a broad and confident first launch.

19.3 Case Study C: Bank with a Wide Network — Edge AI for ATMs

Situation. Thousands of ATMs spread across areas with poor connections; repeated skimming losses and high vendor patrol costs.

Action. Pilot 40 machines across location types with embedded edge devices: fascia anomaly detection + transaction log fusion + sensors. Keys to unexpected success in the initial plan: simulated attack tests (the team installed dummy skimmers) that forced realistic calibration,

and heavy investments in fleet management — zero-touch provisioning and OTA — before adding intelligence. The ATM event is then connected as a feature to the central fraud engine.

Results. Device deployment is detected within minutes (previously discovered days later or after the victim appears), exposed cards are proactively protected before use, and false alarms are controlled thanks to multi-signal fusion. Lesson: at the edge, victory is determined by fleet operations and signal fusion, not by the most advanced models; and the greatest value is often in the connections between use-cases (physical→cyber), not within a single use case.

19.4 Publicly Disclosed Deployments: Evidence, Not Hero Stories

The following examples use public disclosures from the institutions or their technology partners. The figures are useful signals, but they are not normalized benchmarks or independent audits. The point is to identify repeatable operating patterns, not to rank banks.

DBS — Measure Economic Value Before Scaling the Portfolio

DBS reported that more than 2,000 models across over 430 use cases generated approximately SGD 1 billion in economic value in 2025. Its disclosure also connects scale to a governed data platform, reusable GenAI standards, operating-model redesign, and explicit benefit measurement. DBS Joy, a corporate-banking assistant launched in 2025, was used by more than 20,000 corporate and SME customers and was associated with a 23% rise in satisfaction among users. The transferable lesson is not the model count; it is the management system that links reusable capability, adoption, controls, and economic value.

Source: DBS Group Holdings, Annual Report 2025 — CEO Reflections and CIO Statement (dbs.com/annualreports/2025).

Bank of America — Embed AI in a High-Frequency Trusted Channel

Bank of America reported in March 2026 that Erica had exceeded 3.2 billion interactions since its 2018 launch; 20.6 million users interacted with it nearly 700 million times during 2025. Longitudinal adoption matters more than a spectacular launch. Erica is integrated into routine mobile-banking tasks, delivers proactive guidance, and hands customers to specialists when the task requires a person. The lesson is to begin with frequent, bounded customer jobs, instrument every interaction, and expand only when the channel has earned trust.

Source: Bank of America Newsroom, 'BofA AI and Digital Innovations Fuel 30 Billion Client Interactions,' 10 March 2026.

Morgan Stanley — Treat Evaluation as Product Infrastructure

Morgan Stanley and OpenAI reported adoption above 98% among wealth-management advisor teams. Document access reportedly rose from about 20% to 80% as the assistant improved retrieval across institutional knowledge. The implementation emphasis is instructive: domain experts, citation-backed answers, curated content, structured evaluations, and workflow integration were

treated as core product work. The lesson is that enterprise knowledge assistants win through retrieval quality, evaluation discipline, and change management—not through a generic chat interface.

Source: OpenAI, 'Morgan Stanley Uses AI Evals to Shape the Future of Financial Services' (openai.com/index/morgan-stanley).

BBVA — Move from Licenses to Workflow Redesign

BBVA's public case material reports enterprise-scale deployment to roughly 100,000 employees, monthly active usage above 70%, and about three hours saved per employee per week, with larger gains in selected workflows. The more important evolution is from individual productivity to redesigned credit, engineering, risk, and customer-service workflows. The lesson is to treat broad access as the start of transformation: define governed workspaces, reusable assistants, evaluation, data boundaries, and workflow-level benefit measures.

Source: OpenAI, 'BBVA Puts AI at the Core of Banking' (openai.com/index/bbva). Metrics are partner-reported.

OCBC — Pair Broad Access with Narrow, Measurable Tools

OCBC announced in 2023 that it would provide a generative-AI assistant to 30,000 employees. It also disclosed more than four million AI-supported decisions per day and reported productivity gains of up to 50% for selected tools such as document summarization, coding assistance, transcription, and internal knowledge search. The lesson is to separate a safe general assistant from specialized tools with clearer data, quality, and control boundaries, then measure each workflow rather than claiming one enterprise-wide productivity number.

Source: OCBC Group, 'OCBC Is First Singapore Bank to Roll Out Generative AI Chatbot to All Employees Globally,' 24 October 2023.

BRI — Build a Shared Data and AI Hub for a Distributed Bank

Bank Rakyat Indonesia describes BRIBRAIN as a big-data platform that serves as a central hub for AI solutions across business and operational lines. Public information does not provide a directly comparable portfolio metric, so the architectural signal is more useful than a numerical claim: a bank with a very broad distribution network benefits from common data products, reusable decision services, governed access, and feedback loops that connect channels and operations. The lesson is to make the shared foundation explicit before each business line creates its own stack.

Source: BRI Developers, 'Understanding Technology Disruption: Definition, Impact, and Examples' (developers.bri.co.id/en/node/50808).

19.5 Anti-Patterns: The Eight Most Common Ways to Fail

1. **Perpetual PoC** — a portfolio of demos without a single production pipeline; the cure: definition of done = production + business metrics.
2. **Building a platform cathedral without a congregation** — two years of a platform without a use case; the cure: platforms grow from thin slices of real use cases.

3. **Champion model data leaks** — spectacular lab performance collapses in production; the cure: point-in-time discipline and out-of-time validation.
4. **Ignoring integration and process** — right model, no one uses it; the cure: design user workflows from the start, adopt them as KPIs.
5. **Compliance as the final post** — design completed, then rejected; the cure: risk-compliance sits on the team from day one.
6. **No monitoring** — the model rots silently for months; the cure: four layers of Chapter 15 signals as a condition of release.
7. **Vendor black box for material decisions** — cannot be validated, cannot be accounted for; the cure: auditability requirements in procurement.
8. **Overselling AI inwards** — magical expectations, disappointment, pendulum of antipathy; the cure: experiment-based communication and incremental numbers.

19.6 Patterns of Success: Cross-Case Summary

Business owners who charge metrics, not passive sponsors; The first use case has clear value that creates the foundation; labels and feedback loops were designed from the start; automated evaluation discipline before state of the art; investment in ML engineering and operations on par with that in science; the courage to narrow the scope for the sake of quality; and inter-use-case connections as a secondary source of intentional value.

VISUAL TAKEAWAY

Six Public Cases, Four Reusable Patterns



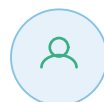
TRUSTED CHANNEL

BofA • Morgan Stanley



SHARED AI HUB

DBS • BRI



WORKFORCE COPILOT

OCBC • BBVA



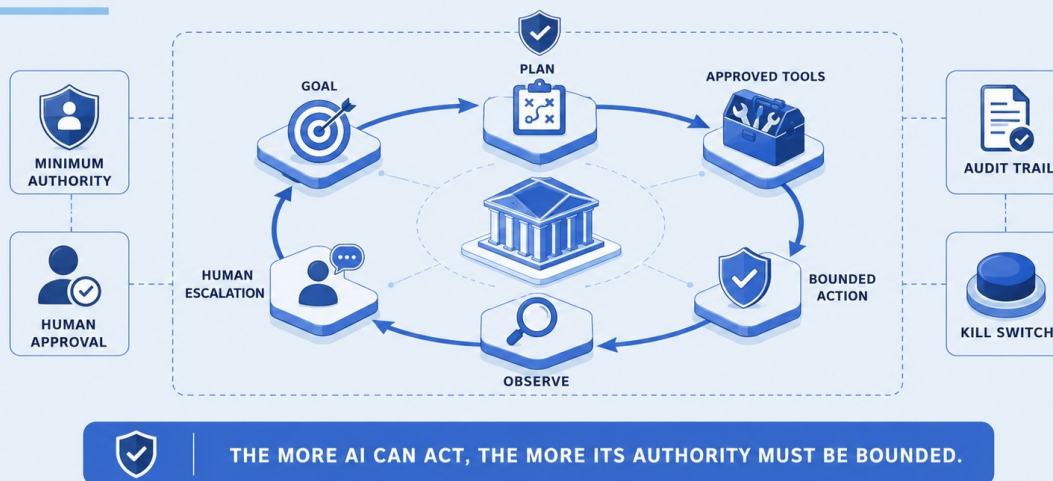
CONTROLLED SCALE

all cases

CHAPTER 20 — Agentic AI and the Future of the AI-Powered Bank

CHAPTER 20

THE FUTURE OF THE AI-POWERED BANK



BUILDING THE AI-POWERED BANK • ENGLISH VISUAL EDITION

Figure 21. The Agentic Banking Control Loop — Future systems may plan and act, but authority must remain bounded and observable.

20.1 From Assistant to Agent: Agentic AI

VISUAL TAKEAWAY Agentic AI should receive the minimum authority required, operate through approved tools, and expose every consequential action.

An agent is an AI system that can pursue a goal across multiple steps, select tools, retrieve information, change system state, check results, and decide what to do next. In banking, that difference is decisive: an assistant proposes, while an agent may act. The opportunity is end-to-end process compression; the new risk is delegated authority. A bank-grade agent must therefore operate inside an explicit decision envelope that defines its identity, tools, data, limits, approval rules, evidence, fallback behavior, and accountable human owner.

20.2 Anatomy of a Bank-Grade Agentic System

A safe agentic workflow is a controlled loop rather than an unrestricted conversation: observe the request and context; plan within policy; retrieve approved evidence; request authorization; execute through registered tools; verify the result against business rules; record the complete trace; and escalate when confidence, authority, or system health falls outside the envelope. Each step must emit evidence that another service or human can challenge.

Control Component	Bank-Grade Requirement
Agent gateway	Authenticates the caller, classifies the task, selects the agent and model, applies rate and cost limits.
Policy decision point	Evaluates purpose, customer consent, data classification, risk tier, transaction amount, and required approval before every consequential action.
Tool registry	Exposes only approved, versioned, typed tools with owners, input schemas, limits, and safe error behavior.
Memory and knowledge	Separates transient task state from governed long-term records; retrieval must preserve provenance and entitlements.
Verification layer	Checks calculations, business rules, citations, state changes, duplicate actions, and post-conditions before completion.
Human control plane	Routes approvals, challenges, exceptions, and emergency stops to named roles with sufficient context.
Audit and replay	Captures prompts, retrieved evidence, policy decisions, tool calls, outputs, approvals, model and tool versions, and final state.

20.3 The Autonomy Ladder

Autonomy should be assigned per task, not per model. The same foundation model may summarize a policy at one level and initiate a payment at another. Banks should begin at the lowest level that produces material value and raise autonomy only after observing stable evidence across representative and adversarial conditions.

Level	Authority	Example	Required Control
L0 — Inform	Answer or summarize; no tool changes state	Knowledge search, policy explanation	Citation, access control, quality evaluation
L1 — Recommend	Prepare a plan or decision for a human	Credit memo draft, fraud investigation pack	Human decision; provenance and limitations visible
L2 — Execute with approval	Prepare tool calls; a human authorizes the consequential step	Card replacement, account-maintenance workflow	Step-up authentication, four-eyes, transaction limits
L3 — Bounded automation	Execute low-risk actions inside pre-approved limits	Case routing, document requests, reversible servicing actions	Policy engine, post-condition checks, sampling, kill switch
L4 — Adaptive orchestration	Coordinate multiple agents and systems across changing conditions	Complex operations in a sandbox or tightly bounded domain	Independent verification, simulation, continuous assurance; exceptional use

20.4 Identity, Authorization, and Human Approval

Every agent requires a workload identity distinct from the employee or customer who requested the task. Credentials should be short-lived and scoped to a named tool, purpose, dataset, customer, amount, and time window. The policy engine must distinguish read, propose, simulate, and commit permissions. Material or irreversible actions require explicit human authorization, while highly sensitive actions may require dual control. Approval screens should show the evidence, proposed state change, residual risk, and a meaningful alternative—not merely an ‘Approve’ button.

20.5 AgentOps: Operating Behavior, Not Just Models

MLOps monitors models; AgentOps must monitor goals, plans, tools, policies, state changes, and recovery. Production evaluation should replay representative tasks whenever a model, prompt, policy, tool, knowledge source, or workflow changes. A release is safe only when the combined system remains within its decision envelope.

Signal Family	Illustrative Measures
Outcome	Task completion, first-pass success, time to resolution, benefit realized
Control	Policy blocks, approval rate, human override quality, limit breaches, unauthorized-tool attempts
Reliability	Tool errors, duplicate actions, partial completion, recovery success, rollback time
Quality	Groundedness, calculation accuracy, plan validity, post-condition verification
Economics	Tokens and tool cost per completed task, human minutes saved, exception cost
Customer and people	Complaint rate, challenge outcomes, user trust, operator workload, escalation quality

20.6 Failure Modes Unique to Agents

Failure Mode	What Happens	Primary Controls
Prompt injection	Untrusted content changes the agent's instructions	Separate instructions from data; sanitize retrieval; tool-level policy checks
Confused deputy	The agent uses legitimate authority for the wrong principal or purpose	Bind identity, customer, purpose, consent, and resource at authorization time
Tool poisoning	A tool, schema, or result is malicious or misleading	Signed registry, version control, output validation, restricted egress
Runaway loop	Repeated planning or tool calls create cost or operational harm	Step, time, cost, and transaction budgets; circuit breakers
Silent partial failure	Some actions succeed while the agent reports completion	Idempotency, transaction boundaries, reconciliation, post-condition verification
Correlated agent error	Multiple agents reinforce the same wrong assumption	Independent verifier, diverse evidence, deterministic business rules
Audit gap	The institution cannot reconstruct why state changed	Immutable event trace with model, policy, data, tool, approval, and result versions

20.7 A Pragmatic Deployment Sequence

Begin with internal L0 and L1 tasks where evidence is available and actions remain human-owned. Move selected servicing workflows to L2 only after evaluation, access control, and approval design are stable. Use L3 for reversible, low-value actions with clear post-conditions and continuous sampling. Treat L4 as an exceptional architecture for bounded domains, not as the default destination. The maturity test is not how independently the agent can act; it is how reliably the institution can constrain, observe, challenge, and recover from that action.

20.8 Generative AI Becomes Infrastructure

Some directions that are almost certain to deepen in the next few years: multimodal models (one model understands text-image-voice-document at a time — changing onboarding, claims, and document operations), small specialized models that are fine-tuned for the domain and run economically on their own infrastructure (addressing data sovereignty and costs), next-generation RAG (knowledge graph + verified retrieval) that increasingly suppresses hallucinations, and AI to build AI — an engineering co-pilot that accelerates the entire cycle in this book. A healthy attitude: treat generative AI as an infrastructure layer that continually changes its engines — your architecture should make changing models as easy as changing tires, because this year's best model will almost certainly not be next year's best.

20.9 Data in Motion: Open Finance and the Ecosystem

The flow of open banking/open finance — the exchange of customer-authorized data via standardized APIs — expands the AI terrain: models that look at the financial picture across institutions provide much better advice and risk assessment, while upping the ante on API consent governance and security. Simultaneously, privacy-preserving collaboration between institutions (federated learning, confidential computing, shared synthetic data) offers a path to fighting cross-bank organized crime without breaching confidentiality — a direction worth investing in now by industry consortia.

20.10 Evolving Threats

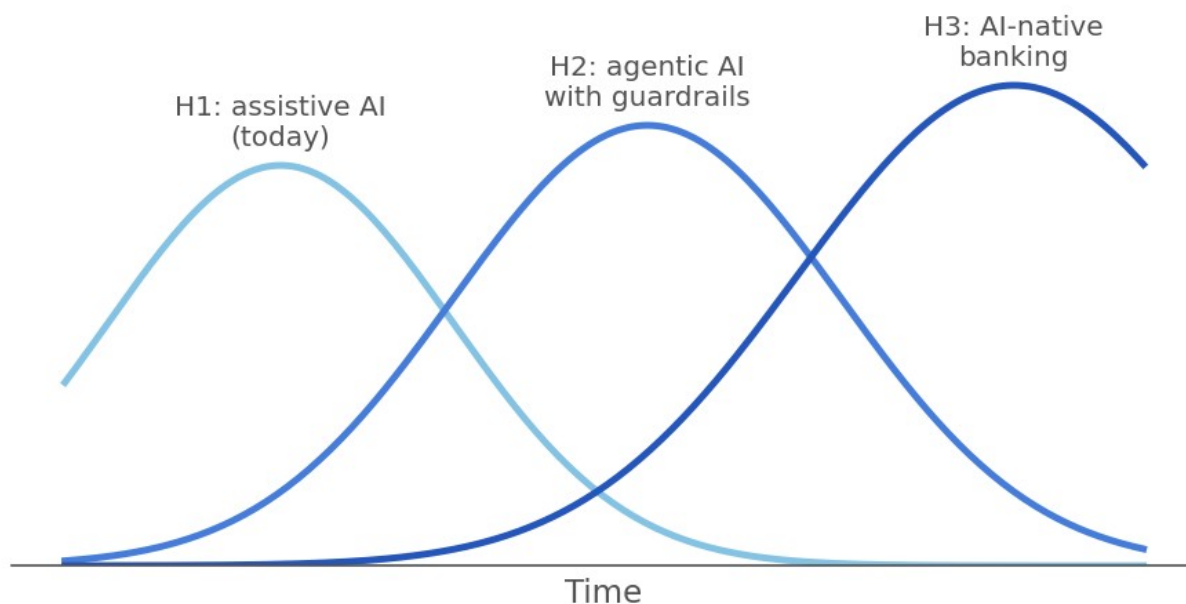
Adversaries use the same technology: voice and facial deepfakes attack identity and authorization verification (the answer: layered liveness detection, device-behavioral signals, and a verification process that is independent of a single modality); LLM-based fraud scales personal and eloquent social engineering; and attacks on AI systems (Chapter 16) become a separate criminal discipline. On a more distant horizon, quantum computing threatens today's cryptography — the gradual migration to quantum-resistant cryptography for long-lived data is work that starts now, not later. Bank defenses must also employ AI: machine learning-based cybersecurity is no longer an option on either side of the fence.

20.11 People at the Center of Smart Banking

A paradox that is both comforting and challenging: the smarter the machine, the more decisive the human. Routine work is automated, and what remains is what is most human — judgment in ambiguity, empathy for customers in difficult moments, ethical accountability, and the question “should we?” which cannot be delegated to the loss function. The banks that win in this era are not the ones with the most models, but the ones that are most successful at weaving machine intelligence

with human wisdom: thoughtfully retraining their people, designing systems that reinforce (not replace) judgment, and maintaining trust — the true currency of banking — in every automated decision.

20.12 Conclusion: The Map Is in Your Hands



From assistive to agentic to AI-native banking

The book opens with an image: a bank branch with nine points of intelligence. Now you hold the complete map — from GPUs in the server room to policies in the committee meeting room. The nine use cases are not the end goal; it's a curriculum: each teaches a different muscle (real-time, edge, conversation, simulation, governance), and banks that master all nine will be able to build their tenth, twentieth, hundredth use case — whatever the future looks like. Start with one, build the foundation right, measure honestly, and take care of the people. Enjoy building your AI-powered bank.

Discussion Questions and Practice

This section provides five questions for each chapter. The questions are designed not to be tested on memorization, but rather to be discussed within the team, used as material for internal workshops, or done as a design exercise. Some questions ask you to take a position and defend it with technical arguments; others require you to design concrete solutions for your own banking context. The best way to use this section: choose the chapter most relevant to your current job, work in cross-functional groups (engineering, business, risk), and write answers in one or two pages so they can be reviewed.

Chapter 1 — The AI-Powered Bank Era

1. Map Westphal Bank's nine use cases to the current state of your bank: which are already running, which are being pioneered, and which are untouched? What patterns do you see from the mapping?
2. Make a counter argument to the claim "AI is a strategic imperative for banks". Under what conditions would a bank be harmed if it rushed to adopt AI?
3. Identify three processes in your bank whose decisions are still completely based on static rules. For each, will machine learning provide significant improvements, or will static rules be sufficient? Explain the criteria.
4. Discuss: why do banking AI transformations fail more often in the organizational adoption stage than in the technical stage? Give real examples from your experience or from public cases.
5. Lay out your team's definition of an "AI-powered bank" in one paragraph, then come up with three measurable indicators that differentiate it from just "a bank that uses some model."

Chapter 2 — Foundations of Machine Learning and Deep Learning

1. For the case of fraud detection with a positive class ratio of 0.1 percent, explain why 99.9 percent accuracy might mean a useless model, and what metrics should be used.
2. A credit scoring model produces a default probability of 0.25 for an applicant. Explain what the numbers mean if the model is well calibrated, and what the business risks are if the calibration is poor but the ranking remains correct.
3. Design simple experiments to detect overfitting in your team's trained models, including data splits, compared metrics, and decision thresholds.
4. When is gradient boosting on tabular data more appropriate than deep learning, and when is it vice versa? Mention at least three determining factors other than accuracy.
5. Explain to a non-technical audience (e.g. product committee) the difference between correlation and causality using one concrete banking example, then explain why uplift modeling is relevant for marketing campaigns.

Chapter 3 — Computing Infrastructure and GPUs

1. Calculate estimated inference requirements for your bank's fraud scoring service: peak transactions per second, target latency, and models used. From there, estimate whether a CPU is sufficient or a GPU is necessary, and explain your assumptions.
2. Discuss the trade-off between building an on-premise GPU cluster and using the cloud for a model training load in a bank that is subject to data residency rules. What factors are most determining in the Indonesian context?
3. Explain why batching improves GPU inference throughput but can hurt p99 latency SLAs, and how dynamic batching tries to balance the two.
4. Create GPU multi-tenancy policies for three competing teams: fraud research, internal LLM development, and data scientist experiments. Include prioritization and isolation mechanisms.
5. What components of the total cost of ownership (TCO) of AI infrastructure are most often overlooked in budget proposals? Name at least five and explain how to estimate them.

Chapter 4 — Data Architecture: Lakehouse, Streaming, and Feature Store

1. Identify one case of training-serving skew that might occur in your bank pipeline, explain the mechanism by which it occurs, and how the feature store prevents it.
2. Explain the concept of point-in-time correctness with the example of the “30-day transaction average” feature for fraud model training, and show how leakage occurs when this principle is violated.
3. Design a four-tier data classification scheme for your bank along with minimum technical controls per tier (encryption, access, masking, retention).
4. When is streaming (e.g. Kafka) really necessary and when is hourly batch sufficient? Give two examples of banking use cases for each and their dividing criteria.
5. Draft a one-page data contract for one of the source tables used by important models in your bank: schema, column semantics, freshness SLA, change rules, and owner.

Chapter 5 — Personalized Branch Banking and Next Best Action

1. Design a framework for determining the “next best action” for customers coming into the branch: action candidates, scoring models, eligibility rules, and inter-offer arbitration logic.
2. Discuss the ethical limits of personalization: when do personalized offers turn into exploitation of customer vulnerabilities? Arrange three concrete guardrails.
3. Explain why A/B experiments are important for the NBA, and design a complete experiment: hypothesis, randomization units, key metrics, safety metrics, and duration.
4. How do you measure the incrementality of an NBA campaign — that sales happen by recommendation, not by chance? Compare the holdout group and uplift modeling approaches.

5. Branch tellers complained the system's recommendations often didn't make sense for customers they knew. Devise a front-liner feedback mechanism that improves the model without opening manipulation gaps.

Chapter 6 — Virtual Financial Advice

1. Compare three virtual financial advice delivery modalities (chat, video advisor, VR/AR) in terms of cost, reach and suitability for Indonesian customer segments.
2. Where is the line between “financial education” and “investment recommendations” according to capital markets regulations, and how does your virtual advice system technically guard that line?
3. Design the escalation flow from virtual advisors to humans: what triggers require escalation, what context is handed over, and how is continuity of conversation maintained?
4. Develop a rubric for evaluating the quality of advice provided by the system: what dimensions are assessed (accuracy, appropriateness of risk profile, completeness of disclosure), who is the assessor, and how many samples per period.
5. Discuss the risks of customers' "over-trust" in smooth-talking virtual advisors. What design interventions can remind you of the limits of a system's capabilities without destroying the experience?

Chapter 7 — Branch Digital Twin

1. Choose one branch operational decision at your bank (e.g. number of tellers per shift) and design how the digital twin helps that decision: the input data, simulation models, and outputs that managers use.
2. Describe the steps to validate that the branch simulator is accurate enough to be trusted, including benchmark historical data and acceptable error tolerances.
3. Discuss the trade-off between simulation fidelity and development costs. How do you decide what level of detail is "good enough" for a particular business question?
4. Data on people's movements in branches is sensitive. Create a privacy policy for the digital twin: what is recorded, at what granularity, how long it is stored, and who has access.
5. Design a “what-if” experiment in the twin to test the new branch layout before a physical renovation: simulated scenarios, success metrics, and how to verify results after a real renovation.

Chapter 8 — Avatar Assisted Banking

1. Map out the ten most frequent customer intents in your bank's digital channels, then classify which ones can safely be fully automated by avatars, which ones require human confirmation, and which ones need to go directly to a human.
2. Design a multi-layered defense strategy against prompt injection on avatars connected to the transaction system: controls in the input, in the model, in tool-calling, and in the output.
3. Explain why LLM system evaluation cannot rely on a single point of accuracy, and develop a minimum evaluation package before the avatar goes into production.

4. Develop avatar persona policies: what to say and what not to say, language style, disclosure that he or she is not human, and behavior when asked outside the scope.
5. Calculate the business threshold: at what level of containment (percentage of conversations completed without a human) and cost per conversation does an avatar become more economical than a contact center, taking into account the cost of errors.

Chapter 9 — IVA for Threat Detection in Branches

1. Compile a list of relevant events detected by smart cameras at your bank branch, sort them based on security value and risk of false detection, then select three for the first phase.
2. Discuss why a “small percentage” false positive rate can cripple a security team operationally. Calculate concrete examples assuming the number of cameras and events per day.
3. Design the alert handling process from detection to incident closure: who verifies, in how many seconds, with what interface, and how decisions are recorded for model retraining.
4. Analysis of the gap between privacy provisions (PDP Law) and video analytics practices: basis of processing, notification to data subjects, retention and access rights. What must be fulfilled before go-live?
5. How do you test the robustness of a detection system to real-world conditions: poor lighting, camera angles, occlusion, and deliberate deception attempts? Design the testing protocol.

Chapter 10 — Data Visualization for Risk Management

1. Select one existing risk dashboard at your bank and critique it based on the principles of this chapter: does it answer decision questions, or does it simply stack graphs? Propose a redesign.
2. Explain the difference between VaR and expected shortfall to non-quantitative managers, including why regulators are shifting to expected shortfall.
3. Design anomaly-based alert mechanisms for risk metrics (not static thresholds): what models to use, how adaptive thresholds are determined, and how to prevent alert fatigue.
4. Discuss the risks of “overly convincing dashboards”: how can neat visualizations hide the uncertainty of the data behind them, and what design elements honestly convey uncertainty?
5. Establish governance controls for the semantic layer of risk metrics: who owns the definition of each metric, how definition changes are enforced, and how consistency between dashboards is maintained.

Chapter 11 — Real-Time Fraud Alerts

1. Select the two fraud typologies that are most detrimental to your bank today, and for each explain the data signals that differentiate it from legitimate behavior and the model that is best suited to detect it.
2. Explain why APP fraud (customer “voluntary” transfers to fraudsters) is much more difficult to detect than account takeover, and what layers of defense it remains possible to build.

3. Design a tiered action policy based on a transaction's risk score: at what score range a transaction passes, is challenged (step-up), held for review, or blocked — and how that threshold is calibrated against the capacity of the review team.
4. Discuss the recall vs false positive dilemma in fraud: how do you translate this technical trade-off into an explicit business decision (fraud loss vs customer attrition vs operational costs)?
5. Fraudsters adapt to your model. Establish an ongoing program to keep the model relevant: drift monitoring, internal red teams, feature updates, and retraining cycles.

Chapter 12 — Automated Investment Advisor

1. Design a risk profile questionnaire that differentiates risk tolerance (psychological) and risk capacity (financial), then explains how the two combine into portfolio placement.
2. Compare three portfolio construction approaches (classic Markowitz, Black-Litterman, risk parity) and explain which one makes the most sense for retail robo-advisors in Indonesia and why.
3. Establish a threshold-based rebalancing rule and analyze its trade-offs against calendar rebalancing: transaction costs, taxes, and risk discipline.
4. Discuss the system's responsibilities when markets fall sharply: what should robo-advisors say and do (and not do) about panicked clients who want to sell everything?
5. Identify points in the robo-advisor flow where capital markets regulations require the involvement or responsibility of a licensed human, and design how the system documents compliance.

Chapter 13 — IVA for ATM Security

1. Create a taxonomy of ATM threats in your bank's operating areas (skimming, card traps, tampering, threats to customers), and for the top three determine the appropriate sensors and detection models.
2. Explain why anti-skimming detection requires a visual baseline per machine, and how the system handles legitimate changes (e.g. replacement of parts) without triggering false alarms.
3. Design an edge architecture for a fleet of thousands of ATMs: what runs on the device, what in the center, how model updates are distributed, and how rollback is performed if updates are problematic.
4. Explain the value of fusing physical signals (cameras, sensors) with cyber signals (transaction logs) — provide one attack scenario that is only detected when the two are combined.
5. Prepare a smart ATM security business case calculation: components of losses prevented, costs per machine, and priority order for which machines are installed first.

Chapter 14 — AI-Powered Bank End-to-End Architecture

1. Draw (sketch) your bank's current AI architecture in five areas (data, model, serving, application, governance), then highlight which areas are most fragile and why.

2. Explain the concept of decision services with explicit fallbacks, and design fallbacks for three services: fraud scoring, NBA, and avatars — what happens when the model is not available?
3. Discuss a strangler pattern for modernizing legacy systems: pick one legacy system at your bank and devise a phased replacement plan without big-bangs.
4. When should a capability be built as a shared platform service and when should it be built locally by the use case team? Arrange the decision criteria so that the platform does not become an obstacle.
5. Design the compute load assignment for your bank's center-region-edge topology: what loads should be at the edge, what at the center, and how model consistency is maintained across all three.

Chapter 15 — MLOps

1. Assess your organization's MLOps maturity: from experimentation to production, which steps are still manual, and which one automation would have the biggest impact if built now?
2. Design four layers of monitoring (infrastructure, data, model, business) for a single production model in your bank, complete with follow-up metrics, thresholds, and owners per layer.
3. Explain the differences between CI, CD, and CT (continuous training) in the context of ML, and discuss when automatic retraining is actually dangerous without validation gates.
4. Put together a short runbook for the “business model metrics are down but technical metrics appear normal” scenario: diagnostic steps from data to delay labels.
5. Discuss the concept of the golden path: what standards are worth standardizing for all ML teams in the bank, and where can teams deviate? How do you enforce it without killing innovation?

Chapter 16 — Security, Risk Management Models, Fairness, and XAI

1. Choose one high-impact model in your bank and map out its attack surface: evasion, poisoning, model theft, and (if LLM-based) prompt injection. Which controls are in place, which are not?
2. Explain the three pillars of the SR 11-7 style risk management model and discuss how to apply them proportionally — so that low-risk models are not burdened by the bureaucracy of high-risk models.
3. Two fairness metrics (e.g. demographic parity and equalized odds) often cannot be met simultaneously. Choose one banking use case and argue which metric is more appropriate and why.
4. Design explanations for three different audiences for the same credit rejection: customers, internal analysts, and auditors/regulators. What is different and what should be consistent?
5. Develop an AI ethics committee's terms of reference (charter): membership, cases to be brought, powers (including stopping the system), and how decisions are documented.

Chapter 17 — Regulations and Compliance

1. Build a mapping matrix from PDP Act obligations to concrete technical controls in your bank's single ML pipeline: processing fundamentals, minimization, retention, data subject rights, and security.

2. Discuss the implications of the "high risk" classification for credit models under the EU AI Act, and what lessons are relevant for Indonesian banks even though they are not directly subject to it.
3. Design a "living evidence pack" for one model: what documents and logs are collected automatically throughout the lifecycle so that the audit doesn't become an impromptu project.
4. Identify points where APU-PPT (anti-money laundering) obligations intersect with ML systems: transaction monitoring, reporting, and explanation of alerts. What are the technical demands?
5. The regulator asked for an explanation of the modeling decision two years ago. Design technical capabilities that enable these requests to be answered: model versioning, feature snapshots, and explanation storage.

Chapter 18 — Strategy and Roadmap

1. Determine where your bank stands on the five levels of AI maturity and lay out three concrete initiatives to move up one level in twelve months — not two levels.
2. Score three candidate use cases in your bank against a foundational value × feasibility × contribution framework, and show how scoring changes (or strengthens) your initial intuition.
3. Discuss the hub-and-spoke organizational model: what functions the hub should handle, what is distributed to the spokes, and how to prevent the hub from becoming a bureaucratic bottleneck.
4. Craft a build vs buy vs partner argument for one specific AI capability, including mandatory contractual clauses so the bank isn't locked into a vendor black box.
5. Design a billing mechanism for ROI realization: how will the benefits promised in the business case be verified quarterly, and what are the consequences if realization falls short of promise?

Chapter 19 — Case Studies and Anti-Patterns

1. Of the eight anti-patterns in this chapter, select the two that are most likely to (or have already) occurred in your organization. What are the early signals, and what interventions can be done now?
2. The avatar case study shows a project that almost failed and then recovered. Analysis: which decision saved him, and could it have been made earlier? What's stopping him?
3. Discuss why "successful pilots" often fail when scaled-up. Name three structural differences between the pilot condition and the full production condition.
4. Establish agreed-upon "kill criteria" for an AI project: what measurable conditions if met will result in the project being terminated without political drama?
5. Write a one-page premortem for an AI project you are or will be running: imagine the project failing miserably two years from now, then write down the most plausible causes and their mitigations today.

Chapter 20 — The Future

1. Discuss agentic AI in the banking context: what transactions or processes in five years might an agent execute autonomously, and what guardrails are absolutely in place before that is permitted?
2. Establish a kill switch policy for autonomous AI systems: who is authorized to suppress, under what conditions, how the system degrades safely, and how those decisions are periodically tested.
3. Voice and video deepfakes attack customer verification processes. Design layered defenses that do not rely on a single detection technique.
4. Discuss the risk of “quantum computing breaking cryptography”: how real will it be for banks in ten years, and what quantum-safe steps make sense to start now without panic?
5. Close with a reflection: from across the book, list the three principles you most want to maintain as technology changes — and explain why they will remain relevant.

Appendix A — AI-Powered Bank Readiness Checklist

Use as an independent assessment per dimension (score 1 = not yet available, 5 = mature-proven). Dimensions with a score ≤ 2 are prerequisites that must be addressed before scaling a use case.

A. Strategy & Governance

- Written AI vision with measurable business goals and executive owners.
- The use case portfolio is prioritized (value $\times$ feasibility $\times$ foundation) and reviewed periodically.
- AI/risk steering committee with real authority (including termination).
- AI policy: principles, prohibition list, approval thresholds per risk level.

B. Data

- Integrated customer identity (entity resolution) across core systems.
- Lakehouse with bronze-silver-gold finish and catalog + lineage.
- Streaming paths for critical domains (payments/channels) with freshness SLAs.
- Online-offline store feature with proven point-in-time correctness.
- Data classification, tokenization of sensitive data since ingest, retention policies in place.
- Automated data quality tests in pipeline + data incident processes.

C. Platform Model & MLOps

- Managed experiment environment with automatic tracking.
- Quoted GPU training capacity; Monitored utilization.
- Registry model with lineage, card model, and approval status.
- Standard deploy path (container + inference server) with shadow/canary/rollback.
- Four-layer monitoring (system, data, model, business-fairness) + runbook.
- Retraining pipeline (scheduled/triggered) with risk-appropriate approval gates.

D. Security & Responsible AI

- Documented AI threat model; periodic red team AI.
- Guardrails LLM + test injection in CI for all generative systems.
- Fairness framework: metrics chosen consciously, tested per group, monitored.
- Explanations are stored with decisions for impact models.
- Independent model validation is functioning at adequate capacity.
- Human supervision is designed and quality measured.

E. Regulation & Privacy

- Regulatory mapping $\rightarrow$ maintained platform control.
- DPIA is mandatory for data-intensive use cases; data subject rights channels are functional.

- Proof of life package per material model (ready to hand over at any time).
- Vendor contracts guarantee auditability, data ownership, and portability.

F. Organization & Adoption

- Hub-and-spoke model with end-to-end use case ownership.
- Key roles filled: ML engineer, data engineer, AI PM, validator.
- Tiered AI literacy program (executive → frontline users).
- Adoption and user feedback are measured as KPIs.
- Benefit realization is billed quarterly with the business owner's signature.

Appendix B — Model Card Template

Contents for each production model; save it in the registry, update each version.

1. **Identity** — model name; version; date; technical owner; business owner; risk level (tiering); approval status and validation date.
2. **Goal & context** — business problem; supported decisions; decision user; coverage limitations; prohibited use.
3. **Data** — source and period; size; label definition; inclusion/exclusion criteria; population representativeness; known quality issues; legal basis for processing.
4. **Methodology** — type of model & rationale for selection; main features; training/validation procedures (division schemes, tuning); dependencies (libraries, base models).
5. **Performance** — key metrics on out-of-time tests; performance per segment/slice matters; calibration; comparison vs baseline/champion; confidence interval.
6. **Fairness & robustness** — tested group; fairness metrics and outcomes; decided trade-offs; stress/adversarial testing and results.
7. **Explainability** — method of explanation; global boost feature; form of explanation available per decision.
8. **Deployment & operations** — serving patterns (real-time/batch/edge); latency budget; fallback policy; active monitoring (metrics & thresholds); retraining trigger; rollback plan.
9. **Limitations & risks** — conditions under which the model is not reliable; critical assumptions; residual risk and its mitigation.
10. **History** — changes between versions; incident and follow-up; next review schedule.

Appendix C — Metrics Catalog per Use Case

Chapter 5 (NBA): conversion uplift vs control; recall@K recommendation; officer adoption rate; incremental revenue per contact; contact limit violation (=0); NPS/branch service time.

Chapter 6 (Virtual advice): sessions completed per month; motion-to-photon latency; CSAT session; advice conversion—transaction vs video call; geographic reach per advisor.

Chapter 7 (Digital twin): simulation-vs-actual difference (<10%); twin-based decisions/quarters; decreased waiting time; VR training competency scores vs legacy methods; avoided design costs.

Chapter 8 (Avatar): containment rate; groundedness/hallucinations (gold test); CSAT; speech-to-response latency p95; 48 hour re-contact; contextual handover success; guardrails incident (=0).

Chapter 9 (IVA branch): precision alert per type; recall of events in planned tests; detection time; alert/operator/hour load; pipeline availability; integrity of evidence.

Chapter 10 (Visualization & predictive): freshness lag per metric; p95 query latency; adoption per persona; forecast vs actual accuracy; lead time early warning; loss of double reconciliation.

Chapter 11 (Fraud): fraud loss (bps); recall@alert-budget; precision alert; customer false positives; detection time; preventable losses; case resolution time; stability per mode.

Chapter 12 (Robo): AUM & funding rate; tracking error vs model portfolio; conformity deviation (=0); customer retention during market corrections; panic-sell rate vs. benchmark; reconciliation error (≈0).

Chapter 13 (IVA ATM): simulated attack recall; skimming detection time; false alarm/machine/month; skimming losses per 1,000 machines; edge device availability; OTA success; cash forecast accuracy.

Across (Chapters 14–18): availability of decision services; p99 latency; idea-to-production time; marginal cost of new use cases; model monitoring coverage; audit findings; benefit realization vs plan; GPU utilization; user adoption.

Appendix D — AI Model Operational Runbook

This appendix presents a ready-to-use runbook for ML operations teams (MLOps/SRE model). The format is intentionally prescriptive: definitions of triggering conditions, diagnostic steps, mitigation actions, escalation paths, and incident closure criteria. Adjust the number threshold to your institution's baseline; the numbers here are a reasonable starting point for the banking environment.

D.1 Principles of Model Operation in Production

A machine learning model in production differs from regular software in one fundamental way: it can “break” without a single line of code changing. The distribution of input data shifts, customer behavior changes, upstream systems modify field formats, or adversaries (fraudsters) adapt. Model operation therefore requires three layers of monitoring running simultaneously:

1. **System health** — whether services are alive, latency within SLA, normal error rate. This layer most closely resembles conventional software operation.
2. **Data health** — whether the input the model receives still resembles the training data: schema, range of values, proportion of missing, feature distribution.
3. **Model health** — whether the model output is still of good quality: score distribution, calibration, and (if labels are available with a lag) actual accuracy metrics.

Golden rule: never wait for labeled accuracy metrics to detect problems. Labels often arrive days to months late (chargeback fraud, credit failure). Early detection relies on proxy signals: input drift, shifts in score distribution, and volume anomalies.

D.2 Runbook 1: Inference Latency Exceeds SLA

Trigger. p99 scoring service latency > 1.5× SLA for 5 consecutive minutes, or p50 > SLA.

Diagnosis (first 10 minute sequence).

1. Check infrastructure dashboard: GPU/CPU utilization, memory, dynamic batching queues. GPU utilization near 100% with long queues → capacity issues. Low utilization but high latency → upstream issues (feature store, networking, serialization).
2. Per-component latency check of distributed tracing: online feature retrieval, pre-processing, core inference, post-processing. Typical experience: 60–70% of latency incidents originate in feature retrieval, not in the model.
3. Check if there were any new deployments (models, applications, infrastructure) in the last 24 hours — correlate the times.
4. Check for traffic spikes: whether the volume of requests increases outside of normal patterns (marketing campaigns, paydays, other channel disruptions that divert traffic).

Multilevel mitigation.

- Level 1: enable manual autoscaling (add inference replicas) if capacity is the cause.
- Level 2: drop expensive optional features (e.g. long-windowed real-time aggregation features) and run the model with a subset of tested features as degradation mode.
- Level 3: switch to a lighter fallback model (e.g. small gradient boosting replacing large ensembles) that is always available warm.
- Level 4: enable rule-based default policies (e.g. all transactions below the nominal threshold pass, above the threshold enter the review queue) — these business decisions must be agreed upon before an incident, not at the time of the incident.

Closing. Latency returns below SLA for 30 minutes; documented root cause; Permanent repair items go into backlog with owners and deadlines.

D.3 Runbook 2: Data Drift Detected

Trigger. The drift score (PSI, KL divergence, or KS test) of key features exceeds the warning (PSI > 0.1) or critical (PSI > 0.25) threshold; or the proportion of missing/invalid features increases > 3× the baseline.

Diagnosis.

1. Identify which features shifted and since when. Simultaneous shift of many features from a single source → suspect upstream system or pipeline changes, not real-world behavior changes.
2. Compare drift across segments: is drift evenly distributed or concentrated in certain segments (regions, products, channels)? Concentrated drift often indicates local events (regional drift, channel disruption).
3. Check the source system change log and data catalog: are there any schema, unit, or encoding changes that were not communicated? Classic case: the nominal field changes from full rupiah to thousands of rupiah, or the category code is updated.
4. Impact score: does the distribution of output scores shift? Input drift without output drift means the model is relatively immune to that feature; still recorded but priority dropped.

Mitigation.

- When the root cause is an upstream data bug: fix at source, backfill impacted features, and assess whether decisions already made during the broken window need remediation (e.g. erroneous transactions are rejected).
- If drift is apparent (change in behavior): schedule retraining with updated data; meanwhile consider adjusting decision thresholds to keep alert/approval volume within operational capacity.
- When drift is extreme and output quality is questionable: reduce automation — increase the portion of decisions transferred to human review until the new model is validated.

Closing. Classified root causes (bugs vs real changes); corrective action completed or scheduled; Monitoring thresholds are reviewed if a false alarm is generated.

D.4 Runbook 3: False Positive Fraud Surge

Trigger. Volume of fraud alerts increased $> 2\times$ daily baseline without an increase in confirmed fraud; or customer complaints "legitimate transactions rejected" spike in complaint channels.

Diagnosis.

1. Take a sample of 30–50 recent alerts, quickly review with a senior fraud analyst: are there any shared patterns (one merchant, one transaction type, one nominal range)?
2. Check the distribution of model scores: is the entire distribution shifting upwards (an indication of a feature/calibration problem) or just certain segments?
3. Correlate with external events: national shopping days, major merchant launches, seasonal behavioral changes (homecoming, payday) that are not yet represented in the training data.
4. Real-time feature health checks: Broken velocity features — for example counters that don't reset — are a classic cause of mass score spikes.

Mitigation.

- Short term: raise temporary alert thresholds for affected segments to keep analyst burden under control; document it as a risk decision with the approval of authorized officials.
- When a feature is broken: fix the pipeline, and consider running the model without the feature (if it has been tested in ablation) rather than with the faulty feature.
- Communication: notify contact centers and digital channels about spikes so that customer response is consistent; set up a fast recovery path for rejected legitimate transactions.

Closing. Alert ratio returns to normal range; affected customers are identified and followed up; A postmortem is scheduled if the customer impact is material.

D.5 Runbook 4: Rollback Model

When to rollback. The new model's output quality metrics deteriorate significantly versus the old model (from shadow/canary comparisons or post-release monitoring); or serious defects are discovered (feature leaks, undetected bias during validation, unexpected behavior in certain segments).

Prerequisites that must always be ready (checked every release, not during an incident).

- Previous versions of models remain registered in the registry with complete artifacts: weights, pre-processing code, feature definitions, and dependency versions.
- Feature compatibility is maintained: older models must still be able to be served by the current feature store. When a new model introduces new features, old features should not be removed immediately — allow a grace period of at least one release cycle.
- Tested rollback procedures: rollback exercises are performed periodically in a staging environment, with time targets (e.g. < 15 minutes from decision to full traffic in legacy models).

Execution steps.

1. The rollback decision is taken by the owner of the MLOps on-call shared service; for regulatory impact models (credit, AML), include risk management representatives according to the authority matrix.
2. Gradually switch traffic (e.g. 100% → 50% → 0% to new model in 10–15 minutes) while monitoring metrics; divert immediately if the defect is critical.
3. Freeze promotion of new models in the registry; mark the version with status and reason.
4. Note the time window when the defect model is active — this is important for impact analysis and possible remediation decisions.

Post-rollback. Blameless postmortem: why the defect passed validation, what evaluation gateways need to be added, and whether the approval process needs to be strengthened. New models should only be re-promoted after the root cause has been fixed and validated.

D.6 Runbook 5: Suspected Adversarial Attack or LLM Abuse

Trigger. Anomalous request patterns to generative/avatar AI services: high volume of few identities, systematically patterned prompts (probing indications), guardrail-block throughput spikes; or reports that the assistant is outputting content that should be blocked.

Diagnosis.

1. Collect indicated samples of conversations/requests (as per privacy and log access policies); classify patterns: repeated jailbreaks, prompt injection via external content, knowledge scraping, or data extraction attempts.
2. Identify source identity/session and trace it across channels; coordinate with the cybersecurity team — treat it as a security incident, not just a model quality issue.
3. Reproduction tests in safe environments: do the same techniques penetrate today's guardrails?

Mitigation.

- Implement rate limiting and authentication step-ups for suspicious sessions/identities.
- Update input/output filters for identified patterns; add the case to the security evaluation suite to prevent regression.
- If there are indications of a data leak: activate the data protection incident procedure (impact assessment, notification obligations according to regulations).

Closing. Documented attack vectors in the internal threat catalogue; guardrail and evaluation updated; If external parties involve malicious intent, follow up according to enforcement/legal policies.

D.7 Model Incident Postmortem Template

Use the following structure, 3–4 pages max, completed within 5 business days of incident closure:

1. **Summary** — one paragraph: what happened, impact (customer, financial, regulatory), duration.
2. **Timeline** — detection, escalation, mitigation, resolution, with timestamps.
3. **Root causes** — the “five whys” technique; distinguish triggers from the latent conditions that enable them.
4. **What went well/badly** — including luck (near misses).
5. **Action items** — specific, proprietary, timed; separate quick fixes from structural fixes; review the resolution in the monthly operations forum.

D.8 Escalation Matrix

Level	Condition	Those Involved	Response Targets
P3	Minor degradation with no customer impact	On-call MLOps	1 business day
P2	SLA breach or degraded output quality	MLOps + product owner	1 hour
P1	Widespread customer impact or systemic incorrect decisions	+ risk management, communications	15 minutes
P0	Suspected regulatory breach / data breach	+ CISO, legal, compliance, senior management	Immediate

Appendix E — Pilot Project Blueprint per Use Case

This appendix presents nine pilot project blueprints, one for each use case on the AI-powered bank vision map. Each blueprint uses a uniform structure: goals, minimum scope, required data, minimum viable architecture, team composition, 12-week plan, success metrics, and go/no-go criteria. The goal is to make the “where to start” discussion concrete: each blueprint is designed to be small enough to last a quarter, tangible enough to measure value, and clean enough to scale to production if successful.

E.1 Pilot Next Best Action (NBA)

Target. Proving that model-based recommendations increase bid conversions in one channel (mobile app) compared to static rules, with valid uplift measurements.

Minimum scope. One channel, 3–5 types of offers (not the entire catalog), one active customer segment (eg monthly mobile users), decisions are calculated daily (not real-time).

Data. 12–24 month transaction history, product ownership, digital channel interactions (clicks, previous offers and responses), basic demographic attributes as per permission. The quality of the offer response history is crucial — if offers have been sent without neatly recording responses, allocate the initial 2–3 weeks to build up the logging.

Minimum architecture. Daily batch scoring on existing data platforms; propensity gradient boosting model per offer; simple decision module (choose the highest score above the threshold, apply eligibility rules and contact frequency); export to existing campaign engine. No need for online store feature yet.

Team. 1 data scientist, 1 data engineer, 1 campaign product owner, part-time support from compliance (eligibility criteria review) and mobile channels.

12 week plan. M1–2: offer definition, response data audit, experimental design (randomized control group 10–20%). M3–5: feature pipeline and model training, offline evaluation. M6–7: integration into campaign engine, end-to-end testing in staging. M8–11: live experiments with controls. M12: uplift analysis, go/no-go decision.

Success metrics. Conversion lift versus control group (realistic initial target: +20–50% relative to static regimen); without a significant increase in complaints/opt-outs.

Go/no-go. Continue if the uplift is statistically and operationally significant; stop/redesign if uplift is unmeasurable — first check the quality of the response data before blaming the model.

E.2 Virtual Financial Advice Pilot

Target. Test whether immersive/visual-rich consulting sessions improve the quality of financial planning conversations and investment product conversions in the affluent segment.

Minimum scope. One scenario (retirement or education planning), one priority branch plus remote session option via browser (pixel streaming) — avoid dependency on headsets in the pilot stage. Interactive Monte Carlo projection visualization at the heart of the experience; Human advisors continue to lead sessions.

Data. Customer profile and financial goals (input during the session), market assumptions for simulation, product catalog with suitability attributes.

Minimum architecture. Interactive 3D/2D visualization applications rendered on GPU servers and streamed to branch devices; projection simulation engine; read-only integration into profile data; session logging for compliance.

Team. 1 3D/frontend application developer, 1 backend/simulation engineer, 1 experience designer, 2–3 branch advisors as test partners, compliance for scripts and disclosure.

12 week plan. M1–3: experience design with advisors, visual prototypes. M4–7: application development and simulation, internal testing. M8–11: real sessions with selected clients (target 30–60 sessions). M12: evaluation.

Success metrics. Session NPS, duration and depth of discussions (number of scenarios explored), progress to product opening, advisor feedback.

Go/no-go. Continue if the advisor considers the tool to be strengthening (not slowing down) sessions and positive conversion indicators; stop if technology dominates but the substance of the advice does not improve.

E.3 Branch Digital Twin Pilot

Target. Building a single-branch digital twin and proving its value on one concrete operational decision: reorganizing the queue/service layout.

Minimum scope. One high-volume branch; geometric model from plan + scan; discrete-event customer flow simulation calibrated with 4–8 weeks of queue and occupancy sensor data.

Data. Branch floor plan/BIM, queue logs (tickets, service times per transaction type), occupancy data (from sensors or structured manual sampling if sensors are not available), staff schedules.

Minimum architecture. Scene 3D (OpenUSD) for visualization; discrete-event simulation engine; calibration-validation notebook that compares simulation output with real observations (p50/p90 lead times per hour).

Team. 1 simulation engineer, 1 3D technician, 1 branch operations analyst, branch manager as partner.

12 week plan. M1–3: geometry and operational data acquisition. M4–6: simulation model and calibration (simulated vs real waiting time error < 15%). M7–9: virtual experiment 3–5 layout/scheduling alternatives. M10–12: physical implementation of best recommendations in the branch, measure before-after.

Success metrics. Model validity (calibration error), real improvement in p90 waiting time after change (target 15–30%), branch manager assessment.

Go/no-go. Continue (replication to other branches, add simulation use cases) if the simulation predictions are proven to be accurate regarding real results; if not, correct the calibration before expansion.

E.4 Pilot Avatar Assisted Banking

Target. Test conversational avatar assistants for 10–15 high-volume service intents at a single service point (branch kiosk or app), with containment and satisfaction as key metrics.

Minimum scope. Low-risk informational and transactional intent (check balance with authentication, product info, branch location/hours, card status, block card with confirmation). Conversational Indonesian, including a mix of informals. Seamless escalation to humans is mandatory from day one.

Data. Legacy contact center/chatbot conversation logs for intent design, product knowledge base for RAG, compliance scripts.

Minimum architecture. ASR pipeline→LLM+RAG→guardrail→TTS→facial animation; function calling to 3–5 service APIs with layered authentication; Conversation dashboard for daily quality reviews.

Team. 1 conversational AI engineer, 1 integration engineer, 1 conversation designer, QA, compliance, contact center representative.

12 week plan. M1–2: guardrail intent and policy design. M3–6: pipeline building and integration, internal red-teaming. M7–8: employee closed test. M9–11: single point of service limited public test. M12: evaluation.

Success metrics. Containment (pilot realistic target: 30–50%), session CSAT, smooth escalation rate (no customers repeating information), zero material guardrail incidents.

Go/no-go. Continue if containment and CSAT are healthy and there are no security incidents/impact hallucinations; if hallucinations/inaccuracies arise, tighten the answer scope to RAG-only before expansion.

E.5 Pilot IVA for Threat Detection in Branches

Target. Proving edge video analytics can produce precise security alerts (not a flood of false alarms) for 2–3 scenarios in one branch.

Minimum scope. Scenarios: abandoned/unused objects, access to restricted areas outside of hours, and loitering in the lobby ATM area. No individual identification; behavior/object based detection only. DPIA was completed before the first camera was analyzed.

Data. Existing 4–8 camera feeds; synthetic data/recorded scenarios demonstrated for training and testing; event labels from security personnel during the test period.

Minimum architecture. GPU edge devices in the branch run the detection-tracking-analytics pipeline; only event metadata and short, time-limited clips are sent to the center; alert console for officers with true/false confirmation button (label for correction).

Team. 1 computer vision engineer, 1 edge/infra engineer, branch security head, DPO/privacy.

12 week plan. M1–2: DPIA, camera mapping, scenario and zone definition. M3–6: development and test with demonstrated scenarios. M7–10: parallel operation (monitored alerts, not yet replacing old procedures). M11–12: precision/recall evaluation per scenario.

Success metrics. Precision alert > 80% (from officer confirmation), adequate recall in controlled tests, alert load per shift within officer capacity.

Go/no-go. Continue if the officer assesses the alert as meaningful and precision achieved; if false alarms are dominant, narrow the zone/scenario and refine the model before adding cameras.

E.6 Risk Data Visualization Pilot

Target. Cut time from data to decision for a single risk domain (e.g. retail portfolio credit risk) with low-latency interactive dashboards and a single layer of agreed metrics.

Minimum scope. One portfolio, 10–15 core metrics with a single definition at the semantic layer, drill-down to account level, daily refresh (intraday to follow), one simple early-warning prediction model with SHAP explanations.

Data. Loan and payment data, collectibility, limits and usage, basic macro data.

Minimum architecture. Existing lakehouse + accelerated query engine for interactive exploration; semantic layer; the BI tool that the institution is already using (do not replace the BI tool in the pilot stage); early-warning model notebook.

Team. 1 analytics engineer, 1 data scientist, 2 risk user representatives (analyst + unit head), metric definition owner from finance/risk.

12 week plan. M1–3: agreement on metric definitions (often the hardest part), building semantic layers. M4–7: dashboard and query latency optimization. M8–10: early-warning model + SHAP. M11–12: guided adoption and evaluation.

Success metrics. Exploration latency (interactive queries < 3–5 seconds at full volume), weekly adoption of target users, reduced number reconciliation between reports, early-warning quality feedback.

Go/no-go. Advanced when risk users move from static reports to self-exploration; when adoption is low, the problem is usually metric definition or data trustworthiness — fix those first, not add features.

E.7 Real-Time Fraud Alert Pilot

Target. Running a modern fraud detection model in the shadow (parallel to the existing system) on one transaction path, proves additional detection with controlled false positives.

Minimum scope. One path (e.g. e-commerce card transactions or digital transfers), real-time scoring within the payment path latency budget, full shadow mode: scores recorded, no blocking.

Data. Labeled transaction history (confirmed + valid fraud) 12–24 months, real-time velocity features, device/channel attributes.

Minimum architecture. Stream processing for velocity features, online feature store, gradient boosting/sequence models, scoring services in the transaction message pipeline (taps, not blocks), score storage for comparative analysis with existing systems and labels that come later.

Team. 1–2 data scientists, 1 streaming engineer, fraud team (analysts and unit heads), payment path engineer.

12 week plan. M1–3: real-time feature pipeline, model training, offline validation. M4–5: shadow integration into transaction pipeline. M6–11: shadow run, weekly review with analysts of decision differences (the model flags those that passed the old system, and vice versa). M12: final analysis with mature labels.

Success metrics. Additional fraud was detected at a false positive rate equal to or lower than the existing system (compare precision-recall curves, not single points); stability of scoring latency.

Go/no-go. Go to gradual active mode (starting from verification step-up, not immediately blocking) when the curve is clearly better; Even if they are equal, whether the operational benefits (clarification, speed of adaptation) still justify continuing.

E.8 Pilot Automated Investment Advisor

Target. Launched a limited-scope robo-advisor: risk profiling, model portfolio recommendations, and rebalancing for select mutual fund products, with full compliance from design.

Minimum scope. 3–5 model portfolios (conservative to aggressive) from mutual funds that have been sold by the institution; validated risk profile questionnaire; recommendation + execution with customer confirmation; proposed rebalancing (not fully automatic) in pilot stage.

Data. Historical product and performance data, customer profiles, daily market prices.

Minimum architecture. Risk profile engine, portfolio construction (documented simple optimization — clarity of methodology is more important than sophistication at this stage), allocation deviation monitoring engine, decision journal for audits, integration into mutual fund transaction systems.

Team. 1 quant/data scientist, 1–2 backend/frontend engineers, investment product management, capital markets compliance (full involvement, not final review).

12 week plan. M1–3: portfolio and questionnaire methodology, compliance review. M4–8: development and testing, including extreme market scenarios. M9–11: launch to limited segments (employees and selected customers). M12: evaluation.

Success metrics. Consistency of compliance (no recommendations outside the risk profile), account activation and funding, customer understanding (brief survey), zero material compliance findings.

Go/no-go. Continue if compliance is clean and adoption trajectory is healthy; remember that a robo-advisor's value grows through scale and trust — evaluation with a longer horizon than other pilots.

E.9 IVA Pilot for ATM Security

Target. Detect two–three threat patterns at ATMs (suspicious loitering, attachment of foreign devices/skimming, vandalism) on 5–10 machines, with alerts that are proven to be actionable.

Minimum scope. Selected ATMs with adequate lighting/cameras; small edge devices per location or per cluster; integration of alerts into existing monitoring centers; without individual identification.

Data. ATM camera feeds, historical incident footage (usually rare — rely on synthetic data and demonstrated scenarios), event logs from monitoring teams.

Minimum architecture. Optimized lightweight detection model for edge (quantization), temporal analytics for loitering, visual change detection on machine fascia for embedded objects, fleet management for OTA model updates, alerts with clip to central console.

Team. 1 computer vision engineer, 1 edge engineer, ATM management team, security.

12 week plan. M1-2: site survey, DPIA, scenario definition. M3-6: model + test with scenario demonstrated on test machine. M7-10: parallel operations at pilot location. M11-12: precision and action-time evaluation.

Success metrics. Precision alerts per scenario, detection-to-action time, edge device reliability (uptime, OTA success), cost per point.

Go/no-go. Proceed to fleet rollout when precision and reliability are achieved and cost per point makes sense at scale projections; consider location prioritization based on risk profile, not equity.

E.10 Ways to Use This Blueprint

Three cross-blueprint notes. First, don't run nine pilots at once — organizational capacity (especially data engineering and compliance) is a real constraint; two–three parallel pilots from different value quadrants is a healthy tempo. Second, go/no-go criteria are established before starting and held; “prolonged without decision” pilots are more expensive than fast-failing pilots. Third, each pilot should leave behind permanent assets — a clean data pipeline, feature definitions in the feature store, DPIA templates, integration patterns — so that even if one pilot stops, the next pilot stands on a higher foundation.

Appendix F — Pattern and Anti-Pattern Catalog

This appendix summarizes design patterns that have proven useful and anti-patterns that repeatedly thwart banking AI initiatives. Each entry is brief: context, problem, solution (or failure and antidote). Use as design review material: before the architecture is approved, check whether relevant patterns have been implemented and relevant anti-patterns have been avoided.

F.1 Design Patterns

Pattern 1 — Feature Store as Contract

Context. The same features are used for training (offline) and presentation (online). **Problem.** Multiple implementations introduce training-serving skew that silently erodes accuracy. **Solution.** Define a feature once as a reversed transformation in the feature store; point-in-time historical value reading training, presentation of current value reading from online stores filled in the same pipeline. Treat feature definitions like public APIs: changes go through review and versioning.

Pattern 2 — Shadow Before Active

Context. The new model will replace the current decision. **Problem.** Offline evaluation does not capture production dynamics (latency, actual data, feedback). **Solution.** Run the new parallel model with no effects (shadow) long enough to collect meaningful comparisons, continue with a canary (small portion of traffic), then rollout fully. For late-label decisions (credit, fraud), plan a shadow window sufficient for label maturity.

Pattern 3 — Layered Graceful Degradation

Context. Real-time decision services must always answer. **Problem.** Component failure (feature store, model) stops business flow. **Solution.** Design an explicit fallback ladder: full model — light model — default rules — each tested and agreed upon by business owners. The system selects the highest rung available in the latency budget.

Pattern 4 — Risk-Appropriate Human-in-the-Loop

Context. Automating customer impact decisions. **Problem.** Full automation is too risky in the beginning; full human review is not scalable. **Solution.** Calibration of automation level against model confidence and decision impact: high confidence + low impact — automatic; additionally — a review queue with sufficient context and explanation for the reviewer. Monitor the quality of human reviews as well (mechanical approval is a silent failure).

Pattern 5 — RAG with Curated Knowledge

Context. LLM assistant answers product/policy questions. **Problem.** LLM without hallucinogenic foundation; Institutional knowledge changes faster than fine-tuning cycles. **Solution.** Limit factual answers to documents drawn from curated and versioned knowledge bases; answers include source references; Out-of-scope questions are directed towards escalation, not answered creatively. Knowledge base content ownership is established (who updates, how often).

Pattern 6 — Edge for Heavy Data, Center for Control

Context. Video/sensor analytics across hundreds of locations. **Problem.** Sending raw video to the center is expensive and poses a privacy risk. **Solution.** Inference at edge locations; only event metadata and relevant short clips rise to the center; the center takes control of the fleet (OTA model updates, configuration, health monitoring). Design edge devices to be physically secure and self-healing.

Pattern 7 — Semantic Layer as a Single Source of Truth

Context. Many dashboards and reports use the same metrics. **Problem.** Metric definitions differ between teams → numbers are inconsistent → trust collapses. **Solution.** Define metrics once in the semantic layer with clear owners; all means of consumption read from there; Changes to definitions are updated and announced.

Pattern 8 — Synthetic Data for Rare Events

Context. Training to detect rare events (robbery, skimming, new types of fraud). **Problem.** Real examples are too few; Waiting for events is not a strategy. **Solution.** Generate synthetic data (3D simulations of digital twins for vision; pattern-based generators for transactions) with controlled variations; always validate the model on existing real data; note the synthetic proportions in the model documentation.

Pattern 9 — Model Registry as the Center of Gravity

Context. Many models, many versions, many environments. **Problem.** "Which models are running in production and what data are they trained on?" missed quickly → audit nightmare. **Solution.** All model artifacts pass through the registry: version, training data (pointer + version), evaluation metrics, approvals, deployment status. Promotion between environments is only through the registry, there are no side channels.

Pattern 10 — End-to-End Latency Budget

Context. Real-time decisions in customer flow. **Problem.** Each component is "fast" but the total is slow. **Solution.** Set a total latency budget from the business SLA, allocate it explicitly per

component (network, feature, inference, post-process), monitor it per component with distributed tracing, and make this budget a design acceptance criterion (new features that exceed the allocation will have to pay with optimization elsewhere).

Pattern 11 — Persistent Control Group

Context. Measuring the value of ongoing decision systems (NBA, pricing). **Problem.** After a full rollout, there is no comparison — unverified value claims and invisible degradation. **Solution.** Maintain a small permanent control group (1–5%) that receives baseline treatment; membership rotation for fairness; report ongoing uplift, not just at launch.

Pattern 12 — Model-Independent Guardrails

Context. Generative systems or automated decisions. **Problem.** Relying for safety on the behavior of the model itself is fragile. **Solution.** Separate layers of input and output checks from the core model: policy filters, format and range validation, source conformance checks (for RAG), action limits (for function calling). Guardrail is tested with its own adversarial suite and updated from incident findings.

F.2 Anti-Pattern

Anti-Pattern 1 — Science Projects without a Production Line

Symptom. Notebook full of good models; neither serves any real decision. **Root.** The research team is separate from engineering; success is defined as accuracy, not value. **Antidote.** Define “done” as a decision to go into production with business metrics; involve engineers from the start; choose the first use case whose path to production is short.

Anti-Pattern 2 — Data Can Wait

Symptom. Modeling begins before data quality and access are addressed; 80% of project time is spent on emergency data corrections. **Antidote.** Data readiness audit as a project start gate; data platform investment is calculated as part of the cost of the first use case, not overhead.

Anti-Pattern 3 — Successful Demo = Production Ready

Symptom. The demonstration amazed the leadership; six months later the production system still doesn't exist because security, scale, and integration have not been touched. **Antidote.** Explicitly differentiate prototype from product in roadmap and budget; load non-functional requirements from design; do not promise a production date from the demo date.

Anti-Pattern 4 — Accuracy as the Only Goal

Symptom. Models are selected solely from offline metrics; lost in production due to latency, cost, resistance drift, or impossibility of explaining to regulators. **Antidote.** Multi-criteria model selection scorecard: accuracy, latency, inference cost, operational complexity, explainability — weighted according to use case.

Anti-Pattern 5 — Compliance as the Final Checkpoint

Symptom. The system is completed, then submitted to compliance; rejected or major overhaul. **Antidote.** Compliance and risk management sat in the design from week one; regulatory requirements are translated into technical acceptance criteria in the backlog.

Anti-Pattern 6 — One Vendor for Everything, or Complete Vendor Allergy

Symptom. Two extremes: total lock-in that paralyzes negotiations and flexibility; or build everything yourself including commodity ones, burning scarce talent on problems that are already solved. **Antidote.** Build-vs-buy map per layer: buy/move fast on commodities (infrastructure, common tools), build on differentiators (institution data signature features, decision integration); maintain portability at expensive exit points (data formats, model interfaces).

Anti-Pattern 7 — The Alert Nobody Reads

Symptom. The monitoring system sends hundreds of alerts; everything is ignored; a real incident was missed. **Antidote.** Each alert must have a clear action and owner; alerts without action are relegated to dashboard metrics; review the monthly alert noise and prune it ruthlessly.

Anti-Pattern 8 — The Model Changes but Its Environment Does Not

Symptom. New model released; contact center scripts, staff training materials, and threshold expectations in business units still refer to old-style behavior; operational chaos. **Antidote.** Model releases are treated as product changes: communication checklists and child artifact updates are part of the definition of complete.

Anti-Pattern 9 — Measure What's Easy, Not What's Important

Symptom. Full report of number of models and experiments; quiet about uplift, savings, or loss reduction. **Antidote.** Each use case has one or two primary business metrics agreed at the start with a valid measurement method (control/comparison); activity metrics are just a complement.

Anti-Pattern 10 — Creepy Personalization

Symptom. Offerings feel too knowing (“You just looked elsewhere for X...”); customers are uncomfortable, trust is down, complaints are up. **Antidote.** Test “fairness” in personalization design (whether customers can understand why this offer arises from their relationship with the bank); respect preferences and frequency of contact; transparent on the basis of personalization.

Anti-Pattern 11 — Eternal Pilot

Symptom. The pilot was “successful” but was never scaled or stopped; hanging around for years absorbing care. **Antidote.** The decision date and go/no-go criteria are set before the pilot begins; governance forums that force decisions; pilot maintenance budgets that are intentionally inconvenient to the status quo.

Anti-Pattern 12 — Talent Bought, Context Not Built

Symptom. Hiring data scientists is expensive; they get frustrated by complicated data access and unclear processes, and leave. **Antidote.** Clear the “paved road” (secure self-service data access, standardized work environment, clear release path) as a prerequisite for large-scale hiring; pair newcomers with banking domain veterans.

Appendix G — Technical FAQs

A collection of the most frequently asked questions from engineers and architects when building banking AI capabilities, with concise answers and relevant chapter references.

G.1 Data and Features

Q: Should we build the feature store before the first use case? A: Not for the first batch of use cases — start with a clean, versioned feature pipeline. Build a feature store when the first real-time use case arrives or when a feature starts to be used across teams; migrating a neat pipeline to a feature store is much easier than tidying up the mess. (Chapter 4)

Q: How long a data history is required for credit/fraud/NBA models? A: Rule of thumb: enough to capture full seasonality plus label dynamics — 12–24 months for NBA and fraud; for credit, ideally cover at least one varying economic cycle, and remember that default labels take 12+ months to mature. More important than length: consistency of field definitions throughout history. (Chapters 4, 11)

Q: How to handle the extreme class imbalance in fraud (0.1% positive)? A: Don't rely on accuracy; use PR-AUC and precision-recall curves on operational thresholds. Techniques: class weighting, careful negative undersampling (keep negatives "hard"), and a focus on engineering velocity/graph features that separate classes better than any sampling tricks. Evaluation is always on the original distribution. (Chapters 2, 11)

Q: What is point-in-time correctness and why is it important? A: When building training data, each feature should have the value it would have known at the historical decision moment — not its current value. Breaking it means future leaks: models look great offline then suck in production. A feature store with time-travel join support solves this systematically. (Chapter 4)

Q: When is synthetic data worth using? A: When target events are rare (physical security, new types of fraud), when privacy prevents using original data for development/testing, and when testing resilience in extreme scenarios. Always final validation on real data; never evaluate a model only on synthetic data from the same generator as the training data. (Chapters 4, 7, 9)

G.2 Modeling

Q: When does deep learning beat gradient boosting for tabular data? A: Less common than expected. Gradient boosting remains a strong baseline for pure tabular. Deep learning excels when there is a sequential (order of transactions), relational (entity graph), or multimodal (text+tabular) structure that can be exploited. Measure the difference against the cost of operational complexity. (Chapter 2)

Q: One big model for all, or many small models per segment? A: Start with one model with segment features; break into segments only if there is evidence of a different fundamental pattern that one model does not capture, and you can afford the operational load of N models (monitoring, retraining, drift — all times N). (Chapters 5, 11)

Q: How to choose a decision threshold? A: Not the default 0.5. Derived from decision economics: false positive vs false negative costs and operational capacity (how many alerts can be reviewed per day). Map the precision-recall/ROC curve to business units (loss dollars, analyst hours) then select operating points with the business owner, and review them periodically. (Chapters 2, 11)

Q: Fine-tuning LLM or just RAG? A: For changing factual knowledge (products, policies): RAG, almost always. Fine-tuning for style, output format, and understanding domain terms — not for injecting facts. A combination of the two is common: light fine-tune for behavior, RAG for content. (Chapters 2, 8)

Q: How to systematically evaluate an LLM system? A: Build a layered evaluation suite: (1) curated test cases with reference answers for factual accuracy; (2) adversarial suite for security (jailbreak, injection); (3) rubric-based evaluation (can be assisted by LLM assessors with human calibration) for quality; (4) production metrics (escalations, corrections, satisfaction). Run on each prompt/model change — a prompt change is a behavior change. (Chapters 8, 15, 16)

G.3 Infrastructure and Operations

Q: On-premise or cloud GPU? A: Typical banking patterns: burst training and experimentation in the cloud (elasticity), stable volume inference and most sensitive data on-premise/colocation (cost per unit and data sovereignty), with an architecture that maintains portability. Calculate TCO at realistic utilization — on-premise GPUs idle 70% of the time cost more than cloud. (Chapter 3)

Q: How to increase low inference GPU utilization? A: Check order: (1) dynamic batching in serving (Triton) — biggest improvement for many-small-request traffic; (2) consolidation of several small models onto a single GPU (MIG or multi-model serving); (3) model optimization (quantization, TensorRT) to make per-request cheaper; (4) autoscaling following daily traffic patterns. (Chapters 3, 15)

Q: How often should models be retrained? A: Not a blind calendar schedule, but a trigger: drift past a threshold, degradation of a labeled metric, or business change (new product/policy). Many teams use a hybrid: basic schedule (monthly/quarterly at domain speed — fraud is faster than credit) plus emergency triggers. Automate the retraining pipeline up to the evaluation gate; Promotion decisions remain subject to approval. (Chapter 15)

Q: What should be in decision logging for an audit? A: For each decision: model version identity and configuration, input feature values (or reconstructable point-in-time pointers), raw scores and final decisions, active rules/overrides, timestamps, and (for impact decisions) resulting explanations. Store within the relevant domain's regulatory retention horizon. Test your “decision replay” capabilities periodically. (Chapters 15, 16, 17)

Q: How to manage the ballooning costs of LLM inference? A: (1) Route by complexity: small models for the majority of simple requests, large models only for difficult ones; (2) cache answers to repeated questions; (3) limit the length of context — don't send the entire history if a summary will suffice; (4) quantization and serving optimization for self-hosted models; (5) monitor cost per conversation as a first-class metric. (Chapters 3, 8)

G.4 Security, Risk and Compliance

Q: How to prevent prompt injection in assistants that read external content? A: Layered defense, not a single filter: strictly separate system instructions from retrieved content (source delimitation and tagging); limit the assistant's action capabilities (function calling with narrow whitelists and confirmations for sensitive actions); filter known injection pattern; and test continuously with an adversarial suite that grows from real findings. Assume the filter will leak — the action limit is the final net. (Chapters 8, 16)

Q: Regulators are asking for credit models to be explainable — should they use simple models? A: Not always. The options are: (a) inherent-interpretable models (scorecard, GAM) when the performance difference is small; (b) complex model + post-hoc explanation (SHAP) with validation that the explanation is stable and faithful; (c) hybrid — complex model for filtering, final decision on interpreted model. Required: reasons for rejection that are meaningful to the customer and defensible methodology documentation. Discuss approaches to the function of risk models since design. (Chapters 16, 17)

Q: How to test model fairness? A: Define protected groups relevant to the Indonesian legal context and agency policy; measure fairness metrics appropriate to the use case (approval rate parity, false positive rate equality — realize these metrics can conflict with each other and the choice between them is a policy decision); also test for proxy discrimination (neutral features that are strongly correlated with protected attributes); document options and mitigations. Repeat each retraining. (Chapter 16)

Q: What are required in a DPIA for a video analytics system? A: Description of processing and its purpose; legal basis; assessment of need and proportionality (why video analytics, why not less invasive means); risks to data subjects and their mitigation (no individual identification, minimal retention, limited access, edge security); transparency mechanisms (signage, privacy policy); and regular review plans. Involve DPO before final design. (Chapters 9, 13, 17)

Q: Who should approve the production ride model? A: Tiered according to risk: low impact model sufficient product owner approval + peer technical review; The model impacts material customer decisions (credit, fraud blocking, AML) through independent validation (a separate risk model function from the developer) and approval according to the institution's model risk management framework. Define the model risk level matrix upfront so that it is not debated on a case-by-case basis. (Chapters 16, 17)

G.5 Organization and Strategy

Q: Centralized AI teams or embedded in business units? A: Typical evolution: starting centralized (density of expertise, standards established), then hub-and-spoke — platforms and standards at the center, use case squads embedded in business units. Purely centralized away from the domain; Purely decentralized duplicates the foundation and disintegrates standards. (Chapter 18)

Q: How to create a business case for a platform (not a use case)? A: Don't sell the platform alone — it's hard to value. Tie the platform investment to the first use case that pays for it, then show the marginal cost curve: the second and third use cases are faster and cheaper because of the platform. The “idea-to-production time” metric per use case is the most persuasive evidence. (Chapter 18)

Q: How long is it realistic to take from zero to the first AI use case in production? A: With fairly clean data and organizational support: 6–9 months for the first batch of use cases including minimum foundation; 9–15 months if real-time or touching highly regulated decisions. Claiming much faster usually comes at the expense of things that will be billed later (security, validation, operations). (Chapter 18)

Q: How to avoid dependence on one or two key people? A: Soft-enforced engineering standards: code and pipelines in a shared repository with reviews, design decision documentation (ADR), service ownership rotation, and operational runbooks (Appendix D) that allow anyone on-call. Bus factor is a real operational risk — treat it as such. (Chapters 15, 18)

About This Book

This book was conceived as an independent engineering guide inspired by the "Building the AI-Powered Bank" (NVIDIA) public vision map. All architectural descriptions, practices, illustrative figures, and composite case studies are written for professional educational purposes; it is not legal, compliance, or investment advice, and is not an official document of any vendor. Apply with adjustments to your institution's context, regulations and risk appetite.

Built-In Implementation Toolkit

Module	Asset	How to Use It
Appendix A	Readiness assessment	Score the institution before funding the roadmap
Appendix B	Model card	Document purpose, data, performance, limitations, ownership, and controls
Appendix C	Metrics catalog	Select business, model, risk, and operational measures per use case
Appendix D	Operational runbook	Respond to latency, drift, false positives, rollback, and attacks
Appendix E	Pilot blueprints	Design a bounded pilot for each of the nine use cases
Appendix F	Patterns and anti-patterns	Reuse proven design choices and recognize recurring failure modes
Appendix G	Technical FAQ	Resolve common architecture, governance, and delivery questions

Updates and author resources: rominur.com

Glossary

Key terms and abbreviations used throughout this book are defined below.

A/B testing - A controlled experiment compares two variants (e.g. with/without a model) in random groups to measure incremental impact.

ABAC (Attribute-Based Access Control) - Access control based on a combination of user attributes, data, and goals, not just roles.

Account takeover (ATO) - The takeover of a customer's account by a fraudster, usually via stolen credentials, SIM swap, or malware.

Adversarial example - Input that is subtly engineered to cause the model to misclassify.

Agentic AI - An AI system that receives a goal and then plans and executes a series of actions using tools/systems to achieve it.

Alert fatigue - Operator fatigue due to a flood of low-precision alerts, resulting in alerts being ignored.

AML/APU-PPT - Anti-Money Laundering / Indonesia's anti-money-laundering and counter-terrorism-financing compliance regime.

Anomaly detection - A technique for flagging patterns that deviate from normality without the need for labeled examples of malicious classes.

API gateway - Managed gateway for API traffic: authentication, quotas, versions, auditing.

ASR (Automatic Speech Recognition) - Technology converts speech to text.

AUC (Area Under the ROC Curve) - A measure of a model's ability to rank positive cases over negative; 0.5 = random, 1.0 = perfect.

Autoencoder - A neural network that compresses then reconstructs input; reconstruction error is used for anomaly detection.

Backpropagation - The algorithm calculates the loss gradient against the network weights for training.

Batch inference - Scheduled bulk prediction over a data set (as opposed to real-time).

Behavioral biometrics - Unique user behavioral signals (typing cadence, touch patterns) for continuous verification.

Black-Litterman - A portfolio construction method that combines market equilibrium with a Bayesian view of investors.

Canary deployment - Release a new version to a small portion of traffic while monitoring before a full release.

Case management - Investigation (fraud/AML) workflow system from alert to decision, with evidence and audit.

CDC (Change Data Capture) - A technique for capturing source database changes as a stream of events.

Champion-challenger - A pattern of running a challenger model side by side with a champion model for continuous comparison.

Chargeback - Refund of disputed card transactions; source of late fraud labels.

Chunking - Chunking a document into pieces for indexing in a RAG system.

CLV (Customer Lifetime Value) - Estimated total value of customer relationships over time.

CNN (Convolutional Neural Network) - Convolution-based deep learning architecture, dominant for images.

Containment rate - The percentage of service interactions completed by an automated system without a human agent.

Contextual bandit - An online learning algorithm that selects actions based on context while balancing exploration-exploitation.

Counterfactual explanation - An explanation in the form of "what needs to change for the decision to change."

CUDA - A parallel computing platform for GPU programming that is the foundation of the deep learning ecosystem.

Data contract - A formal schema, quality, and SLA agreement between data producers and consumers.

Data drift - A shift in the distribution of production input data relative to training data.

Data lake / warehouse / lakehouse - Analytical storage paradigms: lake (raw-flexible), warehouse (structured-managed), lakehouse (a combination of both via open table format).

Data lineage - The lineage of data origins and transformation from source to consumption.

Data mesh - Principles of data organization: ownership per domain, data as a product, self-service platform, federated governance.

Data poisoning - Attacks poison the training/feedback data so that the model behaves incorrectly.

Deepfake - Synthetic media (face/voice) produced by generative AI that imitates real people.

DeepStream - A GPU-accelerated video analytics pipeline framework (NVIDIA ecosystem).

Differential privacy - A mathematical framework limits the leakage of individual information from aggregate/model results via calibrated noise.

Digital human - Realistic avatars with voices and expressions driven by conversational AI.

Digital twin - A synchronized virtual replica of a physical asset or process that uses real or simulated data for testing, training, and optimization.

DPIA - Data protection impact assessment for high-risk processing.

Drift (prediction/concept) - A shift in the distribution of model outputs or input-label relationships over time.

Dynamic batching - Combining inference requests that arrive close together for GPU efficiency.

EAD/PD/LGD - Credit risk parameters: Exposure at Default, Probability of Default, Loss Given Default.

Edge computing - Performing computing near data sources (branches, ATMs) instead of at the center.

Embedding - A dense vector representation of an object (word, customer, product) that captures similarities in meaning.

Ensemble - Combine multiple models for better decisions than a single model.

Entity resolution - Brings together different records that reference the same entity (one customer across systems).

Event-driven architecture - An architecture in which the system communicates via events via a bus/broker.

Exactly-once semantics - Guarantees processing events exactly once even if there is a failure.

Explainability (XAI) - The ability to explain the reasons for a model's output to humans.

False positive/negative - Falsely flagged a positive (false alarm) / failed to flag a positive case (miss).

Feature store - A platform that defines, computes, and serves consistent features for training (offline) and inference (online).

Federated learning - Training models across data locations without moving raw data; only model updates are interchanged.

Fine-tuning - Advanced training of pre-trained models on domain/task-specific data.

FinOps - Cloud/infrastructure cost governance discipline (visibility, accountability, optimization).

Fraud ring - Coordinated syndicate; appears as a dense structure in the entity graph.

Function calling / tool use - LLM capabilities request execution of structured functions/APIs.

GAN (Generative Adversarial Network) - Two-network generative architecture (generator vs discriminator); one of the foundations of synthetic media.

GNN (Graph Neural Network) - Neural network over graph data; learning node representations of relational structures.

GPU (Graphics Processing Unit) - Massively parallel processor; the ultimate engine of modern AI training and inference.

Gradient boosting - An ensemble of sequential trees where each tree corrects the previous error (XGBoost, LightGBM).

Guardrails - A policy layer of input/output filters and behavior constraints for LLM systems.

Hallucinations - LLM output that is glib but not based on facts/sources.

HBM (High Bandwidth Memory) - Very high bandwidth memory on data center GPUs.

Human-in-the-loop - Design that places humans as deciders/reviewers in the AI decision flow.

Inference - The process of a model generating predictions on new input (the opposite of training).

Inference server - A multi-model serving service with batching, scheduling, versioning, and operational metrics; NVIDIA Triton is one example.

IVA (Intelligent Video Analytics) - AI-based video analytics: detection, tracking and understanding behavior from cameras.

Jetson - Family of embedded AI computing modules for the edge (NVIDIA ecosystem).

Kafka - Distributed event logging platform; data streaming backbone.

KS (Kolmogorov-Smirnov) - A statistic that measures the separation between score distributions; it is commonly used in credit scoring and drift testing.

KYC/e-KYC - Know Your Customer; customer identity verification (electronic).

Label leakage - Leaks future information/labels to features, resulting in false lab performance.

Lakehouse - A data architecture that combines the flexibility of a data lake with the management and query features of a data warehouse.

Latency (p95/p99) - Response time at the 95th/99th percentile; a measure of the worst experience commonly used as SLO.

LLM (Large Language Model) - A large-scale transformer language model with the ability to understand and generate text.

Loitering - Staying in a zone beyond a threshold of time; security analytics lock pattern.

LoRA - Parameter-efficient fine-tuning technique with low-rank adaptation matrix.

MIG (Multi-Instance GPU) - Partition a single physical GPU into multiple isolated instances.

MLOps - The engineering practice of operating a model lifecycle: CI/CD/CT, monitoring, governance.

Model card - Model profile standard document: purpose, data, performance, limitations, usage limits.

Model registry - A centralized catalog of model artifacts along with version, lineage, and approval status.

Model Risk Management (MRM) - Banking model risk governance framework: inventory, independent validation, monitoring.

Money mule - An intermediary account for moving the proceeds of crime.

Monte Carlo simulation - Random sampling-based simulation to estimate the distribution of returns (VaR, portfolio projection).

MOT / MOTA / IDF1 - Multi-Object Tracking and its tracking quality metrics.

NBA (Next Best Action) - Recommended next best action per customer per moment.

NLU / NLP - Natural Language Understanding/Processing; natural language understanding and processing.

NPL (Non-Performing Loan) - Problematic credit; core portfolio quality metrics.

OCR - Optical Character Recognition; read text from images/documents.

Omniverse / OpenUSD - 3D collaboration-simulation platform and Universal Scene Description open scene format.

Online/offline features - Millisecond-latency rendered features for inference vs historical features for training.

OTA (Over-the-Air) - Remote software/model updates to a fleet of devices.

Overfitting - The model memorizes the training data and fails to generalize.

PCI DSS - The payment card industry data security standard.

Point-in-time correctness - Guarantee training features reflect exactly the information available at the time of the event.

Pose estimation - Detects body points for movement analysis.

PR-AUC - Area under the precision-recall curve; a key metric for highly imbalanced classification problems such as fraud detection.

Precision / Recall - Accuracy of positive marking / completeness of positive case capture.

Prompt engineering - Designing instructions and context to direct LLM behavior.

Prompt injection - The attack inserts malicious instructions (directly or via captured content) to hijack LLM.

PSI (Population Stability Index) - A measure of shifts in distribution; standard drift score/feature monitoring.

RAG (Retrieval-Augmented Generation) - A pattern of injecting relevant document snippets from search results into the prompt so that LLM answers are source-based.

RAPIDS / cuDF / cuML - GPU-accelerated analytics and ML ecosystem.

Re-identification (ReID) - Recognizing the same object/person across frames/cameras based on appearance.

Rebalancing - Returns portfolio weights to target (calendar based or deviation threshold).

Recommender system - The system recommends relevant items; general pipeline retrieval—ranking.

Reinforcement learning (RL) - Learning action policies from rewards through interaction with the environment.

Retention (data) - Data retention and deletion policy.

Robo-advisor - Algorithm-based, automated investment advice/management service.

ROC curve - TPR vs FPR trade-off curve across thresholds.

RTSP - IP camera video streaming protocol.

Runbook - Standard incident/alarm response procedures.

Scorecard - A highly explained, points-based traditional credit scoring model.

Semantic layer - A centralized metric definition layer on top of the data for BI/analytics consistency.

Shadow deployment - Run a new model on real traffic without using its decisions, for safe comparison.

SHAP - Shapley value-based per-prediction feature contribution attribution method.

SIM swap - Taking over the victim's cellphone number to hijack the OTP/account.

Skimming - Installing illegal devices at ATMs/EDCs to steal card data and PINs.

SLA / SLO - Service level agreements/targets (e.g. latency, availability, data freshness).

SMOTE - Minority class synthetic oversampling technique (use with caution on transaction data).

SR 11-7 - U.S. supervisory guidance on model risk management issued in 2011; it was superseded in April 2026 by SR 26-2.

SR 26-2 - Revised U.S. interagency guidance on model risk management issued in 2026, emphasizing a risk-based approach proportionate to a bank's model-risk profile, size, and complexity.

Stream processing - Continuous processing of data in motion (Flink, Spark Structured Streaming).

Stress test - Testing the resilience of a portfolio/model under extreme scenarios.

Suitability - Suitability of products/recommendations with customer profiles and needs.

Synthetic data - Artificial data distributed similar to real data for development, testing, and training.

Synthetic identity - A fake identity that combines genuine and fabricated elements for application fraud.

Tensor Core - A GPU unit dedicated to mixed precision matrix multiplication.

TensorRT - GPU optimization and runtime inference toolkit (quantization, kernel fusion).

Throughput - Volume of jobs per unit time (predictions/second).

Tokenization (data) - Replacing sensitive data with meaningless tokens; (NLP) breaks down text into tokens.

Training-serving skew - Feature/logic mismatch between training and production; removed by feature store.

Transformer - Self-attention-based architecture that dominates NLP, speech, vision, and sequencing.

TTS (Text-to-Speech) - Natural voice synthesis from text.

Uplift modeling - Modeling the causal effects of actions (the difference in outcomes given vs. no treatment).

VaR (Value at Risk) - Estimated maximum loss at a given confidence level and horizon.

Vector database - Similarity search database over embeddings for RAG/semantic search.

Velocity feature - An aggregate feature of frequency/value per time window (the core of real-time fraud detection).

VMS / PSIM - Video Management System / Physical Security Information Management; physical security platform.

Zero trust - A security model with no implicit trust: continuous verification for every access.

Research Notes and Source Map

This map identifies the principal evidence families used by each chapter. It is a navigation aid rather than a substitute for the complete bibliography. Public deployment metrics are organization- or partner-reported unless stated otherwise and should not be interpreted as normalized benchmarks.

Chapter	Primary Focus	Principal Evidence
1	Strategy and banking AI	NVIDIA, Building the AI-Powered Bank; FSB (2017); OJK, Tata Kelola Kecerdasan Artifiisial Perbankan Indonesia (2025).
2	ML, deep learning, GenAI	Bishop (2006); Goodfellow et al. (2016); Vaswani et al. (2017); Lewis et al. (2020); Murphy (2022).
3	GPU, cloud, edge	NVIDIA CUDA, TensorRT, Triton, and DeepStream documentation; PCI DSS 4.0.1.
4	Data and governance	Kleppmann (2017); Kafka, Flink, Iceberg, and Delta Lake documentation; BCBS 239; Indonesia PDP Law.
5	Next best action	DBS Annual Reports 2024–2025; Guo et al. (2017); OJK consumer-protection requirements.
6	Virtual financial advice	NIST AI RMF; OECD AI Principles; OJK consumer protection; relevant investment-suitability rules.
7	Digital twins	Alliance for OpenUSD documentation; NVIDIA Omniverse and accelerated-computing references where applicable.
8	Avatar and conversation	Lewis et al. (2020); NIST GenAI Profile; Morgan Stanley and OCBC public case disclosures.
9	Branch video analytics	NVIDIA DeepStream documentation; Indonesia PDP Law; OJK technology and cyber-resilience requirements.
10	Risk visualization	BCBS 239; OJK digital-maturity guidance; institutional data-governance and semantic-layer controls.
11	Fraud detection	Breiman (2001); Chen & Guestrin (2016); Hamilton (2020); PPATK risk assessments; Bank of America public disclosures.
12	Automated investment advice	Markowitz (1952); Black & Litterman (1992); OECD AI Principles; suitability and consumer-protection rules.
13	ATM video analytics	NVIDIA DeepStream and edge-inference documentation; PCI DSS 4.0.1; cyber and physical-security controls.
14	End-to-end architecture	Kleppmann (2017); Beyer et al. (2016); BCBS 239; OJK technology-operation requirements.
15	MLOps and reliability	Sculley et al. (2015); Beyer et al. (2016); NIST AI RMF; ISO/IEC 42001.
16	Security and responsible AI	NIST AI RMF and GenAI Profile; ISO/IEC 23894, 42001, and 42005; OECD AI Principles.
17	Regulation and compliance	OJK 11/POJK.03/2022, SEOJK 29/2022, SEOJK 24/2023, OJK AI Governance 2025, PADK OJK 1/2026; EU AI Act; GDPR.
18	Strategy and execution	DBS Annual Report 2025; NIST AI RMF; ISO/IEC 42001; evidence-based benefit-realization practices.
19	Cases and patterns	DBS, Bank of America, OCBC, BRI, Morgan Stanley/OpenAI, and BBVA/OpenAI public disclosures cited in the chapter.
20	Agentic AI and future	IMF Fintech Note 2026/004; Cambridge Centre for Alternative Finance, Global AI in Financial Services Report 2026; DBS Annual Report 2025; NIST AI RMF; OJK AI Governance 2025.

Bibliography

This bibliography consolidates the principal research, standards, regulations, and official technical documentation referenced or recommended throughout the book. Online materials were verified on 12 August 2026. Readers should always consult the latest official version of regulatory and product documentation.

Foundational AI, Machine Learning, and Quantitative Methods

- Bishop, C. M. (2006). *Pattern Recognition and Machine Learning*. Springer.
- Black, F., & Litterman, R. (1992). Global portfolio optimization. *Financial Analysts Journal*, 48(5), 28-43.
<https://doi.org/10.2469/faj.v48.n5.28>
- Breiman, L. (2001). Random forests. *Machine Learning*, 45, 5-32. <https://doi.org/10.1023/A:1010933404324>
- Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785-794.
<https://doi.org/10.1145/2939672.2939785>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press. <https://www.deeplearningbook.org/>
- Guo, C., Pleiss, G., Sun, Y., & Weinberger, K. Q. (2017). On calibration of modern neural networks. *Proceedings of the 34th International Conference on Machine Learning*, 1321-1330.
<https://proceedings.mlr.press/v70/guo17a.html>
- Hamilton, W. L. (2020). *Graph Representation Learning*. Morgan & Claypool.
https://www.cs.mcgill.ca/~wlh/grl_book/
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The Elements of Statistical Learning* (2nd ed.). Springer.
<https://doi.org/10.1007/978-0-387-84858-7>
- Lewis, P., Perez, E., Piktus, A., et al. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459-9474.
<https://proceedings.neurips.cc/paper/2020/hash/6b493230205f780e1bc26945df7481e5-Abstract.html>
- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Advances in Neural Information Processing Systems*, 30. <https://proceedings.neurips.cc/paper/7062-a-unified-approach-to-interpreting-model-predictions>
- Markowitz, H. (1952). Portfolio selection. *The Journal of Finance*, 7(1), 77-91. <https://doi.org/10.1111/j.1540-6261.1952.tb01525.x>
- Murphy, K. P. (2022). *Probabilistic Machine Learning: An Introduction*. MIT Press. <https://probml.github.io/pml-book/book1.html>
- Sculley, D., Holt, G., Golovin, D., et al. (2015). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*, 28, 2503-2511. <https://research.google/pubs/hidden-technical-debt-in-machine-learning-systems/>
- Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998-6008. <https://proceedings.neurips.cc/paper/7181-attention-is-all-you-need>

Data Platforms, MLOps, Reliability, and Accelerated Computing

- Apache Software Foundation. (2026). Apache Flink Documentation. <https://flink.apache.org/documentation/>
- Apache Software Foundation. (2026). Apache Iceberg Documentation. <https://iceberg.apache.org/docs/latest/>
- Apache Software Foundation. (2026). Apache Kafka Documentation. <https://kafka.apache.org/documentation/>
- Beyer, B., Jones, C., Petoff, J., & Murphy, N. R. (Eds.). (2016). Site Reliability Engineering: How Google Runs Production Systems. O'Reilly Media. <https://sre.google/sre-book/table-of-contents/>
- Delta Lake Project. (2026). Delta Lake Documentation. <https://docs.delta.io/>
- Kleppmann, M. (2017). Designing Data-Intensive Applications. O'Reilly Media.
- NVIDIA Corporation. (2026). CUDA C++ Programming Guide. <https://docs.nvidia.com/cuda/cuda-c-programming-guide/>
- NVIDIA Corporation. (2026). DeepStream SDK Developer Guide. <https://docs.nvidia.com/metropolis/deepstream/dev-guide/>
- NVIDIA Corporation. (2026). NVIDIA TensorRT Documentation. <https://docs.nvidia.com/deeplearning/tensorrt/latest/>
- NVIDIA Corporation. (2026). NVIDIA Triton Inference Server User Guide. <https://docs.nvidia.com/deeplearning/triton-inference-server/user-guide/docs/>
- NVIDIA Corporation. (n.d.). Building the AI-Powered Bank. <https://www.nvidia.com/en-us/lp/industries/ai-powered-bank-ebook/>
- OpenUSD Working Group. (2026). Alliance for OpenUSD Documentation. <https://openusd.org/release/>
- PCI Security Standards Council. (2024). Payment Card Industry Data Security Standard: Requirements and Testing Procedures, Version 4.0.1. https://www.pcisecuritystandards.org/document_library/

Responsible AI, Model Risk, and International Standards

- Basel Committee on Banking Supervision. (2013). Principles for Effective Risk Data Aggregation and Risk Reporting (BCBS 239). Bank for International Settlements. <https://www.bis.org/publ/bcbs239.htm>
- Board of Governors of the Federal Reserve System, Office of the Comptroller of the Currency, & Federal Deposit Insurance Corporation. (2026). Revised Guidance on Model Risk Management (SR 26-2). <https://www.federalreserve.gov/supervisionreg/srletters/SR2602.htm>
- Board of Governors of the Federal Reserve System & Office of the Comptroller of the Currency. (2011). Supervisory Guidance on Model Risk Management (SR 11-7; superseded in 2026 by SR 26-2). <https://www.federalreserve.gov/boarddocs/srletters/2011/sr1107.pdf>
- European Parliament & Council of the European Union. (2016). Regulation (EU) 2016/679 (General Data Protection Regulation). <https://eur-lex.europa.eu/eli/reg/2016/679/oj/eng>
- European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Financial Stability Board. (2017). Artificial Intelligence and Machine Learning in Financial Services: Market Developments and Financial Stability Implications. <https://www.fsb.org/2017/11/artificial-intelligence-and-machine-learning-in-financial-service/>

- International Organization for Standardization. (2023). ISO/IEC 23894:2023 - Information technology - Artificial intelligence - Guidance on risk management. <https://www.iso.org/standard/77304.html>
- International Organization for Standardization. (2023). ISO/IEC 42001:2023 - Information technology - Artificial intelligence - Management system. <https://www.iso.org/standard/42001>
- International Organization for Standardization. (2025). ISO/IEC 42005:2025 - Artificial intelligence - AI system impact assessment. <https://www.iso.org/standard/42005>
- National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1. <https://doi.org/10.6028/NIST.AI.100-1>
- National Institute of Standards and Technology. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1. <https://doi.org/10.6028/NIST.AI.600-1>
- Organisation for Economic Co-operation and Development. (2024). OECD AI Principles. <https://oecd.ai/en/ai-principles>

Indonesian Banking, Data Protection, and Financial-Crime Sources

- Bank Indonesia. (2019). Indonesia Payment Systems Blueprint 2025. <https://www.bi.go.id/en/publikasi/kajian/Documents/Indonesia-Payment-Systems-Blueprint-2025.pdf>
- Otoritas Jasa Keuangan. (2022). Peraturan OJK Nomor 11/POJK.03/2022 tentang Penyelenggaraan Teknologi Informasi oleh Bank Umum. <https://ojk.go.id/id/regulasi/Pages/Penyelenggaraan-Teknologi-Informasi-Oleh-Bank-Umum.aspx>
- Otoritas Jasa Keuangan. (2022). SEOJK Nomor 29/SEOJK.03/2022 tentang Ketahanan dan Keamanan Siber bagi Bank Umum.
- Otoritas Jasa Keuangan. (2023). SEOJK Nomor 24/SEOJK.03/2023 tentang Penilaian Tingkat Maturitas Digital Bank Umum.
- Otoritas Jasa Keuangan. (2023). Peraturan OJK Nomor 22 Tahun 2023 tentang Pelindungan Konsumen dan Masyarakat di Sektor Jasa Keuangan.
- Otoritas Jasa Keuangan. (2025). Tata Kelola Kecerdasan Artifisial Perbankan Indonesia. <https://ojk.go.id/id/Publikasi/Roadmap-dan-Pedoman/Perbankan/Pages/Tata-Kelola-Kecerdasan-Artifisial-Perbankan-Indonesia.aspx>
- Otoritas Jasa Keuangan. (2026). Peraturan Anggota Dewan Komisiner OJK Nomor 1 Tahun 2026 tentang Penyelenggaraan Teknologi Informasi oleh Bank Umum.
- Pusat Pelaporan dan Analisis Transaksi Keuangan. (2021). Penilaian Risiko Indonesia terhadap Tindak Pidana Pencucian Uang Tahun 2021. <https://www.ppatk.go.id/publikasi/read/150/penilaian-risiko-indonesia-terhadap-tindak-pidana-pencucian-uang-tahun-2021.html>
- Pusat Pelaporan dan Analisis Transaksi Keuangan. (2024). Penilaian Risiko Sektorial Tindak Pidana Pencucian Uang dan Tindak Pidana Pendanaan Terorisme pada Tindak Pidana Siber Tahun 2024. <https://www.ppatk.go.id/publikasi/read/240/>
- Republic of Indonesia. (2022). Law Number 27 of 2022 concerning Personal Data Protection. <https://peraturan.bpk.go.id/Details/229798/uu-no-27-tahun-2022>

Agentic AI and Public Banking Deployment Evidence

- Bank of America. (2026). BofA AI and Digital Innovations Fuel 30 Billion Client Interactions. <https://newsroom.bankofamerica.com/content/newsroom/press-releases/2026/03/bofa-ai-and-digital-innovations-fuel-30-billion-client-interacti.html>
- BBVA & OpenAI. (2026). BBVA Puts AI at the Core of Banking. <https://openai.com/index/bbva/>
- BRI Developers. (2026). Understanding Technology Disruption: Definition, Impact, and Examples. <https://developers.bri.co.id/en/node/50808>
- Cambridge Centre for Alternative Finance. (2026). 2026 Global AI in Financial Services Report. <https://www.jbs.cam.ac.uk/faculty-research/centres/alternative-finance/publications/2026-global-ai-in-financial-services-report/>
- DBS Group Holdings. (2026). Annual Report 2025: CEO Reflections and CIO Statement. <https://www.dbs.com/annualreports/2025/>
- International Monetary Fund. (2026). Agentic Artificial Intelligence and the Future of Payments. Fintech Note 2026/004. <https://www.elibrary.imf.org/view/journals/068/2026/004/article-A001-en.xml>
- Morgan Stanley & OpenAI. (2025). Morgan Stanley Uses AI Evals to Shape the Future of Financial Services. <https://openai.com/index/morgan-stanley/>
- OCBC Group. (2023). OCBC Is First Singapore Bank to Roll Out Generative AI Chatbot to All Employees Globally. <https://www.ocbc.com/group/media/release/2023/ocbc-is-first-singapore-bank-to-roll-out-generative-ai-chatbot-to-all-employees-globally>

Subject Index

Page numbers exclude the unnumbered front cover. Entries point to substantive discussions rather than every passing mention.

- Agent gateway** 109
- Agentic AI** 108, 120, 147, 154, 158
- AgentOps** 110
- AML / APU-PPT** 11–12, 20, 60, 93, 97, 119, 128, 145, 147–148
- Application programming interface (API)** 8, 11, 43, 50–51, 53, 61, 80–81, 88, 94, 111, 132, 149
- Artificial Intelligence Act (EU)** 93, 95, 119, 154, 156
- ATM security** 10, 25, 74, 117, 135
- Automated investment advisor** 10, 69, 117, 135, 152
- Avatar-assisted banking** 9, 42, 49–50, 52, 100, 115, 132, 148
- Bank Indonesia** 93, 157
- BCBS 239** 154, 156
- Bias and fairness** 22, 39, 68, 72, 84, 86–87, 89–91, 118, 121, 123, 144
- Calibration** 21–22, 32, 38, 46, 48, 65, 75, 84, 89, 99, 104, 113, 155
- Computer vision** 10, 14, 17–18, 46, 97, 133, 136
- CUDA** 23, 28, 148, 154, 156
- Data drift** 16, 126, 148
- Data lakehouse** 11, 30, 33, 36, 38, 61, 80, 98, 114, 121, 133, 148, 150
- Data lineage** 29, 31, 33, 61–62, 80–81, 84, 94, 99, 121, 148, 150
- Data privacy** 11, 26, 30–33, 50–51, 53, 55, 92, 157
- Deep learning** 10, 12, 14, 16, 21–23, 63–64, 89, 113, 142, 148, 155
- Human-in-the-loop** 11, 90, 93, 99, 109, 137, 150
- Inference** 11, 21, 23–26, 28, 31, 33, 38, 51, 55, 61, 156
- Intelligent video analytics (IVA)** 10, 12, 54, 74–76, 100, 116–117, 124, 133, 135, 150
- Kafka** 31, 33, 80–81, 103, 114, 150, 154, 156
- Kill switch** 109, 120
- Large language model (LLM)** 18, 21, 24–25, 27, 42–43, 50–53, 61, 70, 151
- MLOps** 10, 12, 78, 83, 85, 89–90, 97–98, 100, 110, 118, 121, 156
- Model card** 90, 123, 146, 150
- Model registry** 11, 56, 79–80, 138, 150
- Model risk management** 11, 88–90, 95, 97–98, 100, 145, 151–152, 156
- Monitoring and observability** 10–12, 16–17, 25–26, 31, 33, 36, 48, 51, 57, 151
- Next best action** 9, 31, 35–39, 49, 52, 72, 80, 98, 100, 154
- NIST AI RMF** 93, 95, 154
- OJK** 69, 92–93, 95, 154, 157
- Open finance** 111
- Personal Data Protection Law** 32, 92, 95, 116, 154
- Precision and recall** 20–21, 23–25, 27–28, 38, 55, 57, 65–67, 152
- Prompt injection** 19, 21, 50, 53, 87, 90, 110, 115, 118, 128, 144, 151
- Responsible AI** 87, 121, 154, 156
- Retrieval-augmented generation (RAG)** 19, 21, 32–33, 42–43, 49–53, 70–71, 155

Digital twin 10, 44–45, 47, 54–56, 115, 124, 131, 138, 149, 154

Edge AI 25, 28, 104

Explainability 14, 61, 89–90, 99, 118, 123, 140, 149

Feature store 11, 30–31, 33, 38–39, 65–68, 79–80, 97, 153

Fraud detection 10, 12, 14, 63, 113, 134, 151, 153–154

Generative AI 10, 12, 14, 18, 21, 39, 106, 111, 148, 158

GPU 10–12, 17–18, 21, 23–28, 42, 154

SR 26-2 152, 156

Streaming architecture 11, 29, 31, 33, 42–44, 50–51, 55, 66–68, 152

Synthetic data 32–33, 46–48, 55, 57, 111, 133, 135, 138, 142, 152

TensorRT 25, 55, 143, 153–154, 156

Triton Inference Server 25, 143, 150, 154, 156

Zero Trust 153

THE MODEL IS ONLY THE BEGINNING.

Most AI projects do not fail because the model is weak. They fail because data, integration, operations, governance, and accountability were treated as afterthoughts. This visual engineering guide shows how to move AI from an impressive pilot to a safe, measurable, production banking capability.

INSIDE, YOU WILL LEARN TO:

- Design the governed operating stack beneath every AI use case
- Build data, GPU, edge, serving, integration, and MLOps foundations
- Move one use case to production in 90 days and scale it over 12 months
- Control agentic AI with bounded authority, AgentOps, and human accountability
- Measure business value while satisfying model risk, security, and regulation

Romi Nur Ismanto is an AI Research Engineer based in Jakarta with a background in network security and informatics. He writes at the intersection of AI engineering, enterprise architecture, and regulated financial services.

ROMI NUR ISMANTO

AI Research Engineer · rominur.com