

ABC Bank - AI/ML-Powered KYC Uplift Project

Using PCA + t-SNE Clustering to Reduce Outreach Volume and Improve Compliance Efficiency

Company Background

ABC Bank is a large US retail bank that must comply with stringent KYC and AML regulations to ensure customer accounts meet legal and compliance standards. Following recent regulatory updates, ABC Bank identified 35,000 high-risk client records that require immediate review and uplift starting in November 2025. Beyond this initial batch, the bank will need to review and uplift all 10 million client records over time.

Business Problem

Currently, each KYC review requires approximately **7.5 hours** of analyst time, with **80%** of investigations resulting in outreach to clients for additional information or documentation. However, this outreach process is inefficient: only **40%** of contacted clients respond, leading to delays in meeting compliance timelines.

ABC Bank's current KYC uplift process is time-consuming, resource-intensive, and heavily manual. The high percentage of cases requiring outreach—combined with low client response rates—has created operational bottlenecks and increased the risk of non-compliance.

Key challenges include:

Long investigation times, increasing operational costs.

High outreach volume (80% of cases) due to insufficient pre-assessment.

Low client response rate (40%) from outreach through existing phone and email channels.

Limited use of alternative communication channels like SMS or in-app messaging.

The manual, rules-based approach treats all flagged cases equally, wasting resources on low-yield outreach and increasing the risk of regulatory breaches.

Solution Overview – AI-Powered KYC Enhancement

The proposed solution uses **Principal Component Analysis (PCA)** for dimensionality reduction and **t-Distributed Stochastic Neighbor Embedding**

(t-SNE) for visualization, followed by clustering to segment clients into **High, Medium, and Low risk tiers**.

1. Risk Tiering (Clustering / Classification)

High Risk:

- Meets multiple high-risk triggers (e.g., PEP status, high-risk jurisdiction, adverse media, unusual transactions).
- Prioritized for **immediate outreach and full Enhanced Due Diligence (EDD)** by compliance analysts.

Medium Risk:

- Has one or two moderate triggers (e.g., outdated documents, moderate transactional anomalies).
- Managed through **targeted outreach** using the most effective channel (SMS, in-app, secure portal) as recommended by the AI channel-optimization model.

Low Risk:

- Minimal or no red flags, but still due for periodic refresh under regulatory timelines.
- Managed almost entirely through **automated document requests and NLP document handling** with minimal analyst involvement.

2. AI/ML Enablement at Each Tier

High Risk Tier:

- Use AI for **case prioritization** and pre-population of analyst workbenches with aggregated client data.
- Analysts focus on complex judgement calls and escalation to regulators if required.

Medium Risk Tier:

- Use AI-driven **channel selection models** to choose the best way to reach the client.
- Apply **template personalization** to increase engagement and reduce follow-up cycles.

Low Risk Tier:

- Fully leverage **automation** (e.g., auto-reminders, document upload portals).
- NLP-based **document extraction & validation** ensures required fields are complete before analyst review.

Benefits of This Tiered Approach

Faster turnaround for high-priority cases to meet regulatory deadlines.

Optimized resource allocation — senior analysts focus where risk is highest.

Lower operational cost — automation reduces time spent on low-risk cases.

Higher client response rates by using the most effective communication channel for each profile.

What is the cost and time budget for this project?

Cost Budget:

Initial iteration budget: Moderate allocation within compliance operations, targeted at <50% of projected annual savings from reduced analyst hours.

Major cost components:

- Data extraction, cleaning, and transformation (~20%)
- PCA + t-SNE model development, clustering, and SME validation (~25%)
- Visualization & reporting for stakeholder buy-in (~10%)
- Infrastructure (GPU for t-SNE runs, CPU for PCA scoring) (~15%)
- Training, documentation, and integration into KYC workflow (~30%)

Time Budget:

Total project duration: ~10–14 weeks for this iteration

- Data prep: 2–4 weeks
- PCA fit & validation: 1–2 weeks
- t-SNE exploration & SME review: 2–3 weeks
- Clustering & policy mapping: 2–3 weeks
- Deployment & workflow integration: 3–5 weeks

What is the expected ROI for this project?

Operational Savings:

- Reducing outreach rate from 80% to 20% → 60% fewer outreach cases.
- Current outreach process takes an average of **3 analyst hours per case** (on top of initial review).
- On 35,000 cases in this iteration, this yields ~63,000 analyst hours saved.
- At an estimated blended compliance analyst cost of \$60/hour, this equals **~\$3.78M in annualized labor savings** for this batch alone.

Strategic ROI:

- Builds a repeatable, defensible method to scale risk tiering from 35,000 high-priority cases to **10M client records**.
- Frees up analyst capacity for high-complexity cases, reducing regulatory breach risk.
- Supports future AI components (intelligent channel selection, pre-fill & validation, NLP doc handling) to further improve KYC refresh efficiency and client experience.

Regulatory ROI:

- Reduces likelihood of missed SARs or non-compliance penalties by ensuring high-risk cases are prioritized and reviewed on time.

Why does this project need a Cognitive (AI Solution)?

The current manual review process treats **all flagged cases equally**, resulting in excessive manual outreach (80%) and wasted analyst time on low-yield cases. A rules-only approach cannot adapt dynamically to complex patterns in multi-dimensional KYC data.

A **cognitive AI solution** using PCA + clustering (with t-SNE for visual SME validation) can:

- Detect hidden patterns and relationships in high-dimensional KYC data that human analysts or fixed rules would miss.
- Group clients with similar risk profiles, enabling tiered handling strategies.
- Continuously improve as new case outcomes are fed back into the process.
- Produce interpretable outputs (component loadings, cluster characteristics) for regulatory defensibility.

What non-cognitive (non-AI) alternatives are there to solving the current business problem?

Rule-based prioritization

- Use fixed scoring rules based on policy triggers (e.g., PEP = +10, high-risk geo = +7, adverse media = +5).
- Cases above a certain threshold get immediate outreach; others are deprioritized.

Random sampling & manual triage

- Pull a random subset for outreach; adjust sampling rates over time based on yield.

Statistical scoring models (non-AI)

- Use traditional regression or weighted-average scoring without adaptive learning.

For those alternatives, why are they not feasible for this project?

Rule-based:

- Cannot capture complex, non-linear patterns in client behavior or risk factor interactions.
- Requires constant manual updates to rules when regulations or client behavior changes.
- Risks over-prioritizing “obvious” high-risk cases while missing subtle emerging threats.

Random sampling:

- Fails to optimize outreach — wastes analyst hours on low-yield cases.
- Cannot systematically prioritize cases with higher likelihood of confirmed high-risk.

Static statistical models:

- Limited flexibility; performance degrades over time without re-engineering.
- Less effective at handling large volumes of mixed data types (categorical, numeric, binary flags).

If non-cognitive alternatives are feasible, then why are they not being used for this project?

- They **do not meet the primary objective**: reduce outreach from 80% to 20% while maintaining or improving detection of true high-risk cases.
- They provide **lower adaptability** to evolving patterns in transactions, jurisdictions, or risk profiles.
- Regulatory expectations increasingly favor **data-driven, well-validated models** that can explain prioritization logic while continuously improving.

What are the non-cognitive (non-AI) portions of this project that will be used in conjunction with the cognitive components?

- **Policy-based mapping**: Clusters generated by PCA + clustering will be mapped to High/Medium/Low tiers based on SME-reviewed rules.
- **Analyst review process**: All High-risk tier cases will be reviewed by humans; Medium/Low risk cases may have partial human sampling.
- **Compliance governance**: Version control, audit trail creation, and SOP updates remain manual.
- **KYC workflow integration**: Routing logic, notifications, and documentation in the case management system.

Are non-cognitive Automation alternatives possible for this iteration? If so, why aren't they being used for this project iteration?

Yes — robotic process automation (RPA) could automate certain case preparation tasks. However:

- It would **not reduce outreach volume** — the main objective of this iteration.
- RPA would still require humans to decide which cases get outreach.
- RPA is more useful in later phases (e.g., automatically sending outreach to Medium-risk tier cases), after PCA + clustering defines those tiers.

What are the Cognitive objectives for this project?

- **Primary:** Reduce outreach volume from 80% → 20% by accurately clustering and tiering cases.
- **Secondary:** Improve High-risk detection precision without increasing false negatives.
- **Tertiary:** Create a repeatable, scalable tiering model for expansion to 10M client records.

What are the Cognitive outcomes & goals for this project?

- Delivery of **operational PCA + clustering model** for risk tiering, with t-SNE visualizations for SME validation.
- Clear, auditable mapping from clusters to risk tiers aligned with compliance policy.
- Analyst hours saved: $\geq 63,000$ hours in this iteration.
- Operational readiness to add future AI/ML enhancements:
 1. Intelligent channel selection
 2. Automated pre-fill & validation
 3. NLP for document handling

What would the AI project need to be able to successfully do that a non-AI project wouldn't be able to do? In what ways would the AI system need to be better than a non-AI system?

- Identify **multi-dimensional patterns** that correlate with confirmed high-risk outcomes.
- Adapt over time using feedback from analyst overrides and confirmed risk cases.
- Maintain **auditability** by preserving PCA loadings, clustering parameters, and decision mappings.
- Provide **visual evidence** of client grouping for SME and regulator confidence.

Which Pattern(s) of AI are you implementing for this project iteration?

- **Pattern: Predictive Analytics** — clustering for segmentation, risk tier prediction for prioritization.

- **Pattern: Unsupervised Learning** — PCA for dimensionality reduction, clustering for grouping similar cases.
- **Pattern: Explainable AI (XAI)** — feature/component loadings, cluster profiles, and SME-reviewed mappings.

What talent / team resources do you need for this project?

- **Data Scientist (1–2)** — PCA + clustering development, t-SNE visualization, stability testing.
- **Data Engineer (1)** — Data extraction, cleansing, feature engineering.
- **Compliance SMEs (2–3)** — Review clusters, define policy mapping, validate interpretability.
- **Project Manager (1)** — Coordination, timelines, stakeholder communication.
- **IT Systems Integrator (1)** — Deploy models and update case management workflow.

What technology resources do you need for this project?

- **Compute:**
 - GPU (t-SNE runs on sample datasets)
 - CPU cluster for PCA/clustering scoring at scale
- **Software:**
 - Python stack (scikit-learn, cuML/openTSNE, HDBSCAN)
 - Visualization (Plotly, Tableau, or similar)
 - Model registry/version control tools (MLflow, Git)
- **Infrastructure:**
 - Secure cloud or on-prem compute aligned with banking compliance policies
 - Integration APIs for CRM and case management system

What skills do you need for this project iteration?

- Dimensionality reduction and clustering methods (PCA, t-SNE, K-Means, GMM, HDBSCAN)
- Data cleaning and feature engineering for mixed-type data
- Visualization and interpretability techniques
- Model governance and audit preparation
- Financial crime compliance and KYC/AML regulations

What are the project iteration schedule requirements or constraints?

- **Regulatory deadline:** Initial 35,000 flagged high-risk client records must be reviewed and uplifted starting November 2025.
- **Iteration target:** PCA + clustering solution operational within **10–14 weeks** so it's active well before Q2 2026 regulatory timelines.

- **SME review cadence:** Compliance SMEs require 1–2 weeks per review cycle; scheduling delays could push milestones.
- **Data access:** Certain transactional and sanctions screening datasets require internal approvals before use, potentially adding 1–2 weeks.
- **Compute availability:** GPU resources for t-SNE may be shared with other teams, so runs need to be scheduled.

What are the other project constraints that might impact the ability to deliver this iteration?

- **Compliance governance:** All model logic and outputs must be explainable to regulators.
- **Integration complexity:** Model outputs must route seamlessly into the existing case management workflow without disrupting active investigations.
- **Resource limits:** Only one internal data scientist is available; additional contractors must be onboarded quickly.
- **Security:** Strict data handling policies require all processing to occur in bank-approved secure environments.

What are the desired or required performance metrics for the model?

- **Cluster stability:** $\geq 90\%$ consistent membership under bootstrap sampling.
- **Variance retention in PCA:** $\geq 85\%$ with ≤ 30 principal components.
- **Precision for High-risk tier:** $\geq$ current manual process (baseline TBD via historical data).
- **Reduction in outreach:** From 80% $\rightarrow \leq 20\%$ of cases in initial iteration.
- **False negative rate:** $\leq$ baseline to avoid missing true high-risk cases.

What sensitivities are there to false positives or negatives in the case of a binary classifier or inaccurate responses in the case of Generative AI solutions?

- **False negatives (critical):** Missing a true high-risk case could result in regulatory breaches, fines, and reputational damage.
Mitigation: All High-tier cases receive human review; Medium-tier cases receive sampling; Low-tier cases get periodic QA checks.
- **False positives (moderate):** Over-flagging adds analyst workload and erodes efficiency gains, but is less damaging than missing a high-risk case.
Mitigation: Iteratively calibrate tier thresholds and monitor outreach yield rates.

What are the desired or required business KPI performance metrics for this AI project iteration?

- **Outreach reduction:** $\leq 20\%$ outreach rate.
- **Analyst hours saved:** $\geq 63,000$ in initial 35,000-case iteration.
- **Case completion rate:** 100% of High-risk tier cases reviewed within regulatory deadlines.
- **Regulatory audit readiness:** 100% of PCA/clustering decisions traceable and explainable.

What are the desired or required technology KPI performance metrics for this AI project iteration?

- **Model scoring throughput:** Ability to process $\geq 50,000$ cases/day.
- **PCA/clustering execution time:** ≤ 5 minutes per batch of 10,000 cases on production hardware.
- **t-SNE run time:** ≤ 4 hours for 100k sample on GPU (quarterly visual refresh).
- **Data pipeline reliability:** $\geq 99\%$ uptime with automated error alerts.

What, if any, Trustworthy AI Framework will you be using for this project?

- ABC Bank will adopt the **NIST AI Risk Management Framework** and align with **FATF guidance on AI in AML/KYC**.
- Internal AI governance will follow the bank's Responsible AI policy covering fairness, transparency, privacy, and human oversight.

If none, how will you ensure consistent application of Trustworthy AI across this project and others?

N/A — *framework in place*. For cross-project consistency, the same model governance templates, SME review protocols, and bias audit processes will be applied.

What potential physical, financial, emotional, environmental, or other harms could be caused by this project? What approaches will you use to mitigate those potential harms?

Financial harm: Missing high-risk clients → fines, penalties.

Mitigation: Human review for High tier, regular QA for others.

Reputational harm: Over-flagging low-risk customers → complaints, attrition.

Mitigation: Outreach calibration, clear customer communication.

Bias harm: Disproportionate flagging by geography, demographics.

Mitigation: Bias audits, SME review of tier distributions.

***How will you know when the AI project is failing to provide adequate results?
How will you handle AI system failures for this iteration?***

Failure indicators:

- Outreach rate reduction goal (< 20%) not met.
- Precision for High tier drops below baseline.
- Drift metrics indicate significant PCA component change without corresponding policy updates.

Response plan:

- Revert to legacy prioritization rules until model retrained.
- Conduct root cause analysis and SME validation before redeployment.

What do you see as the most significant risks for this project that could lead to project failure?

- Incomplete or delayed data access.
- SME bandwidth limitations delaying validation.
- Misalignment between clusters and compliance policy definitions.
- Over- or under-estimation of true risk in initial tiering.

How will you keep a human in the loop or otherwise involved in the AI project operation?

- All High-tier cases reviewed by analysts before final decision.
- Medium-tier: 50% human review in early phase; adjust as confidence grows.
- Low-tier: Random sampling (10%) plus targeted sampling if drift/bias detected.
- All overrides logged and fed back into retraining datasets.

How will you go about identifying and minimizing exposure to informational bias?

- Remove or control for variables highly correlated with protected characteristics.
- Conduct bias audits by geography, customer segment, and product type.
- SME validation of cluster assignments to ensure no unjustified skew.

For this AI project iteration, what laws, regulations, or other compliance might be required? If you don't know, how will you find out?

- BSA (Bank Secrecy Act)
- USA PATRIOT Act
- FinCEN KYC/AML guidance
- OFAC sanctions compliance
- FATF recommendations
- Internal ABC Bank AML/KYC policies

Compliance office will ensure full mapping of model logic to these regulations before production deployment.

What transparency are you going to provide to others about the source(s) of the data used in this AI Project?

- Document and disclose all internal and external data sources used (KYC profiles, transactions, sanctions lists, adverse media feeds).
- Maintain a data source inventory for audit purposes.

What transparency are you going to provide to others about the methods you use to select and filter the data you're using for your AI project?

- Document feature selection criteria, preprocessing steps, and any exclusions (e.g., dropped features due to high missingness or bias risk).
- Keep preprocessing code and config files version-controlled for reproducibility.

What are the requirements for explainable algorithms for this AI project?

- PCA loadings and explained variance report for interpretability.
- Cluster profiling to show key features driving group membership.
- SME-reviewed mapping from clusters to policy-defined tiers.
- Stored decision trail for each case (PCA scores, cluster ID, tier assignment).

Detail the list of data and locations of that data you will need for this iteration of the AI project.

Primary Data Sources (internal):

- **KYC profiles** (CRM system) – personal details, risk flags, onboarding date, doc completeness, PEP status.
- **Transaction history** (Core banking system) – last 2 years, aggregated metrics (velocity, cash intensity, geographies).
- **Sanctions & watchlist screening results** – OFAC, UN, EU, internal lists.
- **Adverse media hits** – internal adverse media screening tool output.
- **Prior outreach history & response rates** – from case management system.

Secondary Data Sources (internal & external):

- External business registries (for corporate clients).
- External identity verification databases.
- Geopolitical risk scoring feeds.

If you are encountering any issues with locating or accessing data, document resolution to these issues.

- **Transaction monitoring system** exports require special compliance sign-off → resolution: secure export pipeline approved by Data Governance.
- **Adverse media data** includes licensing restrictions → resolution: Legal review completed; usage for internal ML approved.

Document the nature of the data you need. What structure is it? Does it have the elements that you need for your AI project iteration?

- **Structure:** Tabular data (mixed numerical, categorical, binary flags).
- **Elements needed:** Complete case identifiers, core KYC profile attributes, aggregated transaction metrics, sanctions/media flags, historical outreach and outcome fields.
- **Status:** Elements available; some categorical features need encoding, some numerics need normalization.

Have you inspected and selected some of the data to make sure it meets your needs?

Yes — sample inspection shows:

- Numeric transaction fields have skew → will require log transformation.
- Some categorical features (e.g., country codes) have inconsistent formats.
- Missingness mostly in secondary features (adverse media score, corporate ownership structure).

What is the current quality of the data you located for your AI project?

- **KYC profiles:** High completeness; occasional outdated doc expiry dates.
- **Transactions:** High quality; rare missing fields due to processing delays.
- **Sanctions & watchlist hits:** Clean; binary flags.
- **Adverse media:** Medium quality; some false positives, will require SME review.

What needs do you have for data preparation, augmentation, enhancement, and transformation?

- Standardize categorical codes (countries, product types).
- Log-transform skewed numeric variables.
- Create derived risk metrics (transaction velocity, geographic diversity score).
- Encode missingness indicators for use in PCA.

What additional, specific needs do you have for training data for your AI project?

- Confirmed “true high-risk” labels from past EDD cases for back-testing precision/recall after clustering.

Select the Appropriate Algorithm and approach to be used for model development.

- **PCA** – reduce dimensionality while retaining $\geq 85\%$ variance.
- **t-SNE** – visualize 2D/3D structure for SME interpretation (not for production scoring).
- **Clustering** – K-Means (fast baseline), GMM (soft membership), HDBSCAN (detects dense pockets).
- **Mapping** – SME-driven mapping from clusters → High/Medium/Low tiers.

Perform data cleansing and preparation operations.

- Outlier handling (winsorization at 1% tails).
- Missing value imputation (numeric: median, categorical: mode).
- Encoding: one-hot for nominal categorical, ordinal for ordered categories.
- Standardization: z-score scaling for PCA compatibility.

Determine approach used to validate the model and ensure that it doesn't overfit or underfit.

- Cluster stability testing via bootstrap sampling.
- Cross-validation of PCA components on stratified subsets.
- Back-testing with holdout set of confirmed high-risk cases.

Evaluate the model and produce evaluation measures.

- **Technical metrics:** Variance retained, silhouette score, Davies–Bouldin index, cluster stability %.
- **Business metrics:** Outreach reduction %, precision/recall for High tier, analyst hours saved.

Detail required approvals or reviews to be conducted before the model can be operationalized in production.

- Compliance SME sign-off on cluster→tier mapping.
- Model Risk Management review for adherence to internal model governance policy.
- IT Security review of deployment infrastructure.

How will this model be operationalized in what mode and in what location(s)?

- **Mode:** Batch scoring nightly; outputs written to case management system.
- **Location:** On-prem secure compute cluster, with GPU node reserved for quarterly t-SNE runs.

What continuous monitoring and management approach and tools will be used for the model in this iteration?

- Monthly drift checks on PCA component scores.
- Quarterly bias audits on tier distributions.
- Outreach yield rate monitoring (cases leading to confirmed high risk).
- Model retraining trigger if drift/bias or performance thresholds breached.

What should be done in the next iteration for this AI project?

- Integrate **Intelligent Channel Selection** model to improve response rates from outreach.
- Implement **Automated Pre-Fill & Validation** to reduce client friction.
- Deploy **NLP for Document Handling** to cut analyst review time further.

Perform an iteration post-mortem.

(To be completed after initial deployment)

- **What went well:** TBD
- **What didn't go well:** TBD
- **Improvements for future iterations:** TBD