

Transaction and synthetic identity fraud detection for smaller institutions

OPEN METHODOLOGY · FREE TO RUN, VERIFY, AND ADOPT

IMPLEMENTATION GUIDE · VERSION [VERSION]

What it takes to run this, what it costs, and how to read what comes out

Written for a community bank, a credit union, or a regional lender with no data science team. It assumes one technically capable person and a compliance officer, and it tells you what you can do at three levels of investment, including the level where you write no code at all.

THREE TIERS START AT THE BOTTOM. MOST INSTITUTIONS SHOULD NOT BEGIN ABOVE TIER 1.				
TIER	WHAT YOU NEED	SETUP EFFORT	RUNNING COST	WHAT IT BUYS YOU
TIER 0 A manual review queue	Your own application history in a spreadsheet or a database. No external data, no model, no code.	Two to three days for one person who can write a query.	Review staff time only	Catches the clustering patterns: shared addresses, shared contacts, one identity number under several names.
TIER 1 Scripted scoring, run in a batch	Tier 0 plus a machine that can run Python, and whatever credit file data you already pull at origination.	Three to six weeks part time for one technically capable person.	Compute is negligible. Staff time dominates.	A ranked queue refreshed nightly, with reasons attached to each ranked application.
TIER 2 Scoring inside the origination flow	Tier 1 plus integration with your origination system, and a model risk process your regulator will accept.	Three to six months, and it is mostly not the modelling.	Integration, governance, monitoring	Decisions informed at the point of application rather than the morning after.

The compute is not the expensive part, and that surprises people. These models run on an ordinary CPU with no accelerator, and the software is free. What costs money is the time of the people who work the queue, which is why the cost model in Section 3 is built around review hours rather than around servers.

ASSUMES One technical person Plus a compliance officer. No data scientist.	SOFTWARE COST Zero Python, pandas, scikit-learn, all open source	HARDWARE CPU only No accelerator at any tier
DATA MOVEMENT None Everything runs inside your environment	COMPANIONS Feature specification [LINK] . Read Part 4 first.	VERSION [VERSION] Published [DATE]
LICENSE CC BY 4.0 Adapt it to your institution	ARCHIVED AT [DOI] Permanent identifier	

1 Whether you should do this at all

Four questions. If the answers point the wrong way, the useful outcome of this guide is that you stop here and spend the money on something else.

Q1 Do you originate credit remotely?

Synthetic identity fraud is an origination problem and it travels through remote channels. An institution that opens accounts almost entirely in branch, with identity documents physically inspected, has meaningfully less exposure and should weigh this accordingly. **If your remote origination volume is small and not growing, Tier 0 is probably the correct ceiling.**

Q2 Can you see your own charge off history?

The cost model in Section 3 depends on knowing what a loss costs you. If nobody can pull charge offs by origination channel and vintage, fix that first. It is cheaper, faster, and more useful than anything in this guide, and it is a prerequisite for knowing whether this work paid for itself.

Q3 Is there anybody to work the queue?

A ranked list nobody opens is worse than no list, because it creates a documented record that the institution was alerted and did nothing. **Do not build a queue until somebody's job description includes working it.** That person does not need to be an analyst; they need to be an underwriter or a fraud reviewer with authority to pause an application.

Q4 Will your compliance officer be in the room from the start?

Anything that influences a credit decision brings obligations with it: adverse action reasons, fair lending testing, model risk expectations, and record retention. Those are cheap to design in and expensive to retrofit. If compliance first hears about this when it is built, expect to rebuild it.

NOT LEGAL ADVICE

This guide describes obligations in general terms so that you know what to ask about. It is not legal advice and the author is not your counsel. Every regulatory point in Section 5 must be confirmed with your own compliance officer and, where a credit decision is affected, with counsel.

2

Tier 0: what you can build this week, with no model

Five measurements from the feature specification need nothing but your own application history. They address the structure of the scheme directly and they are the cheapest thing in this document.

A synthetic identity is a real identity number under a fabricated name, and operators run them in batches. That means the same address, telephone number, email, or employer shows up across applicants who have no relationship to each other, and the same identity number shows up under more than one name. **None of that requires a credit bureau to see. It is sitting in your own application table.**

TIER 0 QUEUE, BUILT FROM YOUR OWN RECORDS ONLY

SIGNAL	WHAT TO COMPUTE	HOW TO READ IT
Name multiplicity	Count of distinct names appearing against the same identity number across the trailing <code>[SET, suggested 1095]</code> days.	More than one name against one number needs an explanation. Marriage and transcription errors account for some of it, and not for all of it.
Address multiplicity	Count of distinct applicants sharing a normalised address in the same window. Fix your normalisation rule and write it down, because formatting differences will otherwise fragment the count and hide the cluster.	Multi occupancy is real, so this ranks rather than decides. A single address with eight unrelated applicants in a month is not multi occupancy.
Contact multiplicity	Count of distinct applicants sharing a telephone number, or the local part of an email address, in the same window.	Contact details are reused more than addresses because they cost nothing to hold and are rarely checked across applications.
Application clustering	Count of other applications in the trailing 30 days sharing two or more identity elements with this one, where an element is address, telephone, email local part, or employer.	Two shared elements between otherwise unrelated applicants is difficult to explain innocently. This is usually the highest yield line in a Tier 0 queue.
Income coherence	Ratio of stated income to the median stated income for the same occupation and geography in your own history, with a minimum group size of <code>[SET]</code> .	Benchmarks the field most often inflated against your own book rather than against an external table.

HOW TO RUN TIER 0

Rank applications by how many of the five they trip, break ties by application amount, and put the top `[IN]` per week in front of a reviewer. There is no score, no threshold, and no model. There is a list, and every line on it can be explained to the applicant in a sentence, which matters for Section 5.

Run Tier 0 for at least one full quarter before considering Tier 1. It will tell you your base rate, it will tell you how long a review actually takes, and both of those are inputs to the cost model that you cannot get any other way.

3

Data required, by tier

What each source gives you, whether you probably already have it, and what to do if you do not.

SOURCE	TIER	WHAT IT ENABLES	IF YOU DO NOT HAVE IT
Your application history	0, 1, 2	All five Tier 0 signals, and the reference populations everything else is ranked against.	You do have it. The usual obstacle is that applications are stored per decision rather than in one queryable table. Fix that first.
Your charge off and loss history	0, 1, 2	The loss figures in the cost model, and later the only labels you will have.	Without it you cannot state whether this paid for itself. Treat as a prerequisite rather than an enhancement.
Credit file data you already pull	1, 2	File age, authorised user structure, tradeline velocity, inquiry velocity, limit growth. This is the bulk of the specified feature set.	Most institutions pull this at origination already and discard it after the decision. Retain it. That single change is often the whole gap.
Account performance over time	1, 2	Utilisation trajectory, synchronised drawdown, payment regularity. The measurements that catch a bust out before it completes.	Requires a monthly panel per account rather than a current balance snapshot. Start retaining month end snapshots now, since you cannot reconstruct them later.
Third party identity data	2	Non credit footprint, element age, address occupancy history.	The most expensive input and the least necessary. It improves the specified feature set, and it is not required to get value. Do not let its absence stop you.
Consortium or shared data	2	Cross institution clustering, which is where a scheme cultivated across several lenders becomes visible.	Sector bodies and some core processors offer this. Ask before assuming you cannot get it.

THE ONE ACTION THAT COSTS NOTHING AND MATTERS MOST

Start retaining the credit file data you already purchase, and start taking month end account snapshots. Both are cheap today and impossible to recover retrospectively. An institution that begins retaining now has a usable panel in twelve months. An institution that waits until it has decided to build something starts that twelve month clock on the day it decides.

4

What it costs, and how to work out whether it pays

Fill in your own numbers. The structure is what is being published here; the figures have to be yours, and any figure supplied by a vendor should be treated as a starting point for the same worksheet.

Where the money actually goes

- **Software: zero.** Python, pandas, and scikit-learn carry no licence cost. There is nothing to buy to start.
- **Compute: negligible.** These models train and score on an ordinary CPU. An institution scoring a few thousand applications a month can do it on a machine it already owns, and a modest cloud instance covers a much larger book. No accelerator is required at any tier.
- **Setup staff time: the main Tier 1 cost.** Three to six weeks of part time work for one technically capable person, and most of that is data plumbing rather than modelling.
- **Review staff time: the main ongoing cost, by a wide margin.** Every item you put in the queue consumes a reviewer's attention whether it turns out to be anything or not. This is the number the whole exercise turns on.
- **Data: variable and optional above Tier 1.** Retaining what you already buy costs nothing. Purchasing new identity data is the largest single line item and the last one you should add.
- **Governance: real, and routinely forgotten.** Documentation, validation, fair lending testing, and monitoring. At Tier 2 this can exceed the build.

Monthly cost worksheet

MONTHLY OPERATING COST		
Review hours per month Queue depth multiplied by average minutes per review, from your Tier 0 quarter	A	
Loaded hourly cost of a reviewer Salary plus benefits and overhead, not salary alone	B	
Analyst maintenance hours per month Monitoring, refresh, and investigating drift. Budget more than you expect.	C	
Loaded hourly cost of the analyst	D	
Data costs per month Only what is additional to what you already buy	E	
Compute and hosting per month	F	
Total monthly cost $(A \times B) + (C \times D) + E + F$	COST	

MONTHLY BENEFIT		
Applications reviewed per month Your real capacity, not an aspiration	G	
Share of reviewed applications that turn out to be a genuine problem Measured in your Tier 0 quarter. Do not guess this.	H	
Average loss avoided per case prevented From your own charge off history, net of recoveries	J	
Realisation factor Share of identified cases you actually stop in time. Start at 0.5 and measure it.	K	
Total monthly benefit $G \times H \times J \times K$	BENEFIT	

THREE WAYS THIS WORKSHEET GETS FUDGED

Using salary instead of loaded cost for B and D. It understates the true cost by a large factor and is the most common error in an internal business case.

Guessing H. The share of reviewed applications that turn out to be something is the single most sensitive number in the model, and it is the reason Tier 0 runs for a quarter before anything is built. A number taken from a vendor deck or a published study is a number about somebody else's book.

Setting K to 1.0. Identifying a synthetic identity and stopping it in time are different events. Some cases are found after the drawdown, and those cost the same as not finding them. Measure the gap rather than assuming it away.

IF THE WORKSHEET SAYS NO

Then it says no, and that is a legitimate and useful result. A small institution with low remote origination volume may find the review cost exceeds the avoided loss at every queue depth. Record that, and revisit it if remote volume grows. Nothing in this guide is improved by an institution running it against its own arithmetic.

The section people skip and the one where institutions get into difficulty. A score is a statement about a pattern, not about a person, and everything downstream depends on holding that distinction.

The output is a **rank within a defined reference population**, together with the measurements that drove it. It is not a probability that an applicant is fraudulent, it is not a determination, and it is not evidence. It says that this application resembles a pattern that has historically been worth a look, relative to a population you have specified.

That framing is not lawyerly caution. It changes what you do with the output. A rank triggers a review; the review produces a decision; the decision is made by a person who can state a reason. An institution that collapses those three steps into one will eventually decline somebody it cannot explain declining.

RANK BAND	WHAT IT MEANS	WHAT TO DO
Top [X] %	Resembles the cultivation and clustering patterns strongly, relative to your reference population. Usually several measurements agree.	Hold for manual review before approval. Verify identity through a channel the applicant did not choose. Document what the reviewer found, including when they found nothing.
Next [Y] %	One or two measurements are unusual and the rest are ordinary. Frequently explained by a genuine circumstance such as a recent immigrant with a thin file.	Proceed with standard verification and note the flag on file. Do not add friction here unless your review capacity is genuinely unused.
Remainder	Nothing in the pattern stands out.	No action. Do not record a fraud related note on an application that was not reviewed.

Set the bands from your capacity, not from the scores

Decide how many applications your reviewers can genuinely hold per week, and let that number define the top band. Choosing a score cut off first produces a volume the team cannot absorb, then a backlog, then a quiet decision by reviewers to stop opening the queue. That failure gets blamed on the model afterwards, but it was designed in at this step.

Reasons, not scores, reach the applicant

Where an application is declined or given less favourable terms, the applicant is entitled to the specific principal reasons. "Our model scored you highly" is not a reason. This is the practical

argument for preferring measurements that can be stated in a sentence a customer would understand, and it is why the specification favours them.

Confirm the exact requirements with your compliance officer. In general terms, expect obligations around specific principal reasons for adverse action, additional notice requirements where a consumer report contributed, fair lending testing of outcomes, supervisory expectations for models that influence credit decisions, and retention of the inputs and the reasons for a defined period.

Four things never to do with this output

- **Never decline on a score alone.** The score selects who gets looked at. A person decides, and records why.
- **Never describe an applicant as fraudulent on the basis of a rank.** Not in a file note, not in an email, not to another institution. A pattern resemblance is not a finding, and the note outlives the context.
- **Never share a score outside the institution** without a legal basis your counsel has confirmed. Consortium sharing is legitimate and has its own rules.
- **Never let the queue run unworked.** An unworked alert queue is a documented record that you were notified and did nothing, and it is worse than not having built it.

6

A ninety day rollout to Tier 1

Assumes one technically capable person at roughly half time, a reviewer with authority to pause an application, and a compliance officer available for two conversations.

Weeks 1 to 2 Prepare

Get the data into one place, and get compliance in the room

Assemble applications, decisions, and charge offs into a single queryable table. Confirm what credit file data you already purchase and whether it is retained after the decision. Begin retaining it if not, and begin month end account snapshots.

Hold the first compliance conversation now, framed as: we are building a review queue, it will not make decisions, and here is what the reviewer will see.

Weeks 3 to 4 Tier 0

Build the five signal queue and start working it

Compute the five Tier 0 measurements. Rank, take the top items per week, and put them in front of the reviewer. Record for every item: minutes spent, outcome, and whether the reviewer would have caught it under existing triage.

This is the measurement phase The minutes and the outcome rate are inputs H and A in the cost worksheet. Without a real quarter of them, the business case at week 12 is guesswork.

Weeks 5 to 8 Build

Add the credit file measurements and produce a ranked score

Implement the origination observable families from the feature specification: file genesis, element coherence, tradeline cultivation, and application behaviour. Compute a percentile rank within an explicitly defined reference population and publish the population rule alongside it.

Attach the top contributing measurements to every ranked application. A rank without reasons cannot be reviewed efficiently and cannot support an adverse action.

Weeks 9 to 10Test

Compare against your existing process, blind

Mix flagged applications with matched controls from below the threshold, unlabelled, and have reviewers work them normally. Compare the yield. Fix the acceptance margin in writing before you unblind.

The comparator is your current process Not a published benchmark and not a vendor's figures. If the queue you already run performs as well at the same depth, the correct outcome is to keep what you have.

Week 11Audit

Fair lending and reason quality

Compute flag rates across protected classes and record the result. Confirm that the sex field is not in the feature matrix and that geography is not acting as a proxy. Read twenty sample reason sets aloud and ask whether each could be stated to an applicant.

Week 12Decide

Fill in the worksheet and write the decision down

Complete both halves of the cost model with measured rather than assumed numbers. Adopt, pilot with a capped queue depth, defer pending better data, or stop. Whichever it is, record it and the numbers behind it, so that a successor understands why.

7

How this goes wrong

Six failure modes, each of which has a specific and cheap prevention.

The queue nobody worksThe most common

Built by a technical person, handed to an operations team that was not consulted, opened for two weeks and then not at all.

Prevention: the reviewer works the Tier 0 queue from week 3, before anything is automated. If they will not work five items a week by hand, they will not work fifty from a model.

The threshold set by the modelAlert fatigue

A cut off chosen from a score distribution produces a volume the team cannot absorb, and reviewers begin clearing items without reading them.

Prevention: start from capacity, take that many off the top, and read the threshold off the ranking.

The score treated as a finding
The most damaging

A file note that says suspected synthetic identity, written on the strength of a rank, follows an applicant and gets repeated by people who never saw the underlying reasons.

Prevention: file notes record what the reviewer verified, not what the model ranked. Write that rule down and audit against it.

Training on the wrong label
Quiet and expensive

Synthetic identities that never draw down are never labelled, and those that do are usually recorded as charge offs rather than fraud. A model trained on those labels learns to predict charge off.

Prevention: at Tier 1, rank on the specified measurements rather than fitting a classifier. Add supervision only when you have labels somebody can define in a sentence.

The panel that was never retained
Costs a year

The utilisation and payment behaviour measurements need month end snapshots. Current balance only tells you today, and history cannot be reconstructed.

Prevention: start snapshotting in week 1, before you have decided whether you are building anything.

Nobody owns it after go live
Slow decay

The person who built it moves on. Distributions drift, a data feed changes format, the ranking quietly degrades and nobody notices because nobody is looking.

Prevention: name an owner at week 12, set a re validation date, and track realised outcome rate against the estimate. Divergence is the earliest signal that something upstream changed.

8

What this guide does not give you

Stated plainly, so that an institution knows what it is taking on before it starts.

- **No trained synthetic identity model exists.** The feature specification defines the measurements; it has not been fitted or evaluated against any label. Tier 1 as described here is a ranked queue built from specified measurements, not a validated classifier.
- **No performance figures for synthetic identity detection are available,** from this work or, as far as the author can establish, in a form comparable across institutions. Anyone quoting one should be asked what population it was measured on.

- **The transaction detection results were produced on synthetic datasets.** They demonstrate that the feature construction captures signal across two structurally different environments. They do not establish behaviour on your production data.
- **This guide does not cover integration with any specific origination or core system.** That work is real, it is where most of the Tier 2 timeline goes, and it depends on vendors this document knows nothing about.
- **It is not legal or regulatory advice.** Section 5 tells you what to ask about. Your compliance officer and counsel tell you what applies.
- **It cannot tell you the scheme will not adapt.** Fraud responds to detection. Whatever you build needs a re validation date from the day it goes live.

IF YOU TAKE ONE THING FROM THIS DOCUMENT

Retain the credit file data you already purchase, start taking month end account snapshots, and run the five signal Tier 0 queue by hand for a quarter. That costs almost nothing, it is useful on its own, and it is the only way to obtain the numbers that tell you whether anything further is worth doing.

Run it, and tell me what it cost you

This guide exists to remove acquisition cost and implementation burden as reasons a smaller institution has no detection capability. If a step here took far longer than stated, if a cost line was missing, or if the tier ladder does not match how your institution actually works, that is a defect in this document and the most useful thing you could report.

COMPANION DOCUMENTS

- Feature engineering specification [\[LINK\]](#)
- Reference implementation [\[LINK\]](#)
- Threat model note, synthetic identity [\[LINK\]](#)
- Healthcare validation protocol [\[LINK\]](#)

AVAILABILITY

Methodology and this guide are published without restriction. Executable implementations are available free to verified institutions, confirmed automatically against public registries, with no approval step, no fee, and no size exclusion.

CORRECTIONS

Report a cost line that was wrong, a step that did not work, or a regulatory point stated inaccurately:

[\[CONTACT\]](#)

Ayeni, A. [YEAR]. Implementation guide: synthetic identity and transaction fraud detection for smaller financial institutions, version [VERSION]. [DOI] · Licensed CC BY 4.0.

Educational and technical material only. Nothing here is legal, regulatory, or financial advice. The output described is a prioritisation instrument for human reviewers and does not establish that any applicant has committed fraud.