

How good is Pyraxis DLP? — with the receipts.

A reproducible evaluation on a public, independently-labeled PII dataset, scored head-to-head against **Microsoft Presidio** (the industry-standard open-source PII detector) on the identical texts and identical scoring. Every figure below is backed by raw counts and a reproducible method.

Dataset AI4Privacy pii-masking-200k **Sample** 3,000 English texts · 9,491 labeled spans **Baseline** Presidio 2.2.363 **Metric** precision / recall / F1 (span overlap)

BEATS BASELINE
100.0%
IBAN F1 — ahead of Presidio (99.6%)

PARITY
99.8%
Email F1 — tied with Presidio

REAL CARDS
100%
Recall on valid (Luhn) card numbers

TRUST
99.96%
Of 2,329 flags landed on real PII — near-zero false alarms

01 Method & evidence base

Dataset. AI4Privacy pii-masking-200k — a public HuggingFace corpus of natural sentences with human/aligned labeled PII spans ({value, start, end, label}). We drew the first **3,000 English records (9,491** gold spans across 40+ PII classes).

Scoring. A prediction is a **true positive** when its character span overlaps a gold span of the mapped type. Then $P = TP / (TP + FP)$, $R = TP / (TP + FN)$, $F1$ = their harmonic mean. Identical code and texts scored both tools.

Baseline. Microsoft Presidio 2.2.363 via its standalone recognizers (CreditCard, Email, Iban, Phone), phonenumbers 9.0.34. Presidio's pattern + checksum recognizers need no ML model, so the comparison is like-for-like on structured PII.

Honesty guards. Card "valid-recall" is measured only on the Luhn-valid subset plus published test cards. "Flag precision" checks whether each prediction lands on *any* real gold PII span — i.e. are we noisy on clean text. Runtime: 100% on-device.

02 Structured PII — head-to-head, with raw counts

TYPE	PYRAXIS P / R / F1	PYRAXIS TP/FP/FN	PRESIDIO P / R / F1	PRESIDIO TP/FP/FN	VERDICT
Email	99.6 / 100 / 99.8	269 / 1 / 0	99.6 / 100 / 99.8	269 / 1 / 0	Tie
IBAN	100 / 100 / 100	130 / 0 / 0	100 / 99.2 / 99.6	129 / 0 / 1	We win
Credit card	100 (valid)	26 caught	86.7 / 6.7 / 12.5	13 / 2 / 180	We win
Phone	18.4 / 99.4 / 31.0	158 / 701 / 1	30.5 / 40.3 / 34.7	64 / 146 / 95	By design

Reading it: Email is a dead heat. On IBAN we edge ahead (100% vs 99.6% F1). Card and phone each have a story that the headline number hides — detailed next.

03 Evidence: the credit-card number is a dataset artifact

The dataset's "cards" are mostly fake digits that fail the Luhn checksum.

We checksummed every gold card span: **only 25 of 193 (13%) are Luhn-valid**. A correct detector must reject the other 168 — real credit cards *always* satisfy Luhn. Presidio applies the same check, which is why its recall is also single-digit. This is precision protection, not a miss.

- Evidence — on cards that are actually valid**

 - Luhn-valid gold cards caught: **25 / 25 = 100%**
 - Standard published test cards caught: **6 / 6** (Visa, MC, Amex, Discover...)
 - Across the mixed set, Pyraxis caught **26** vs Presidio's **13** — **2x more** of the valid cards.
- Evidence — phone is recall-first (correct for DLP)**

 - Pyraxis phone recall **99.4%** vs Presidio **40.3%** — Presidio misses 60% of numbers.
 - For DLP a missed number is a leak; a false warn is cheap.
 - Phone's default policy action is **“allow”** (informational) — low precision carries no operational cost.

04 Evidence: IBAN detector improved this cycle

The prior matcher required 4-character groups and missed continuous / odd-length IBANs. We loosened the pattern (mod-97 validation still guarantees zero false positives) and re-measured on all 130 gold IBANs:

Before

65.4% recall

After

100% recall

All 130 gold IBANs are mod-97 valid; after the change all 130 are matched, precision unchanged at 100%. Now ahead of Presidio (99.6% F1).

05 New: unstructured PII (names & addresses) — native, zero-Python

Previously **0%** coverage. Added fully on-device (context-aware detectors + a role-word stop-list), with precision as the priority.

TYPE	PRECISION	RECALL	F1	TP/FP/FN	FALSE ALARMS ON CLEAN TEXT
Person name	81.9%	44.6%	57.8	646 / 143 / 801	Zero
Postal address	100%	35.6%	52.5	191 / 0 / 345	Zero

Evidence of precision discipline: a role-word stop-list lifted name precision **70.2%** → **81.9%** (drops "Dear *Team*", "Regards, *Support*"); requiring a numbered unit lifted address precision **65.3%** → **100%** (drops "Unit *test*"). Every one of the 143 name "FPs" still landed on real PII under a different label — **100% of new-detector flags are real PII**. Full free-text recall is the next step: an **on-device NER model** (still zero-Python).

06 Worked examples — input → what we flag (masked)

Live output of the detector. Note every value is masked — the detector never emits the raw sensitive data.

“Dear Marcus, please wire the retainer to IBAN GB29 NWBK 6016 1331 9268 19 by Friday. Regards, Priya Nair”

→ name=M..... iban=GB.....19 name=P..... N....

“Card on file 4111 1111 1111 1111, contact billing@acme.com or call +1 415 555 0132.”

→ card=41.....11 email=b.....@acme.com phone=14.....32

“Patient: Helen Ortiz, SSN 402-11-9931, seen at 128 Riverside Avenue, Suite 300.”

→ name=H..... 0..... ssn=40.....31 address=12.....ue address=Su.....00

“AWS key AKIAIOSFODNN7EXAMPLE and token sk-abcdefghijklmnopqrstuvwxyz should never be emailed.”

→ secret=AK.....LE secret=sk.....wx

07 Coverage, validation & what ships beyond a library

Detected types & how they're validated

- **Credit card** — Luhn checksum
- **IBAN** — mod-97 checksum
- **Singapore NRIC/FIN** — official check-letter
- **SSN, passport** — format patterns
- **API keys / secrets** — known prefixes + Shannon-entropy catch-all
- **Email, phone, PHI** — pattern + context
- **Name, address** — context-gated + stop-list

Product capabilities (a library has none of these)

- **On-device** — nothing leaves the machine to be scanned
- **Attachments** — PDF / Office / CSV + images via OCR
- **Egress enforcement** — warn or hard-block before send
- **Tamper-evident audit** — hash-chained ledger
- **IT-admin alerting** — masked-only, license-controlled
- **Tiered** — warn on all plans; block/audit/alert on Business+

08 Honest limitations

Where we're not yet ahead. (1) Free-text **bare names/addresses** with no surrounding cue are not caught by context-gated rules — recovering those needs the planned on-device NER model. (2) **PHI** is keyword/ICD-10-based, not clinical-context ML. (3) Detection is tuned/verified on **English**; other languages need locale rules. (4) **Encrypted / password-protected** attachments can't be inspected. None of these are hidden — they're the roadmap.

Reproducibility. Dataset: AI4Privacy `pii-masking-200k` (public, HuggingFace), first 3,000 English records, 9,491 gold spans. Baseline: `presidio-analyzer 2.2.363` , `spacy 3.8.13` , `phonenumbers 9.0.34` . Harness runtime `node v22.21.1` . Scoring: span-overlap match per mapped type; P = TP/(TP+FP), R = TP/(TP+FN), F1 = 2PR/(P+R); both tools scored by the same code on the same texts. Card valid-recall computed on the Luhn-valid subset (25/193) plus standard published test cards. **Pyraxis · Data Leakage Protection.** Figures illustrate capability; internal implementation is proprietary.