



Edge AI-Enabled Screening Tools: A Solution Blueprint for Real-Time Data Management and Operational Efficiency

Solution Blueprint 1.0

Author: Beenish Zia

Contributors: Sarah Canny, Jaime Gregar, Hassnaa Moustafa, Gabriel Silva

Intel® Edge Computing Group – Healthcare and Life Sciences

Forward-Looking Statements Disclaimer

Statements in this document that refer to future plans or expectations are forward-looking statements. These statements are based on current expectations and involve many risks and uncertainties that could cause actual results to differ materially from those expressed or implied in such statements. For more information on the factors that could cause actual results to differ materially, see our most recent earnings release and SEC filings at www.intc.com.

Table of Contents

1. Executive summary

2. Market analysis

3. Technical requirements.

4. Target workflows (HST use cases)

5. Reference architecture

5.1 System-level architecture:

5.2 Component-level architecture

5.2.1 Reference system components

5.2.2 Functional requirements

5.2.3 Non-functional requirements

5.2.4 Inference execution and workload placement

5.2.5 Observability and benchmarking

6. Reference hardware configuration — screening tools

7. Data handling, privacy, and governance

8. Potential Edge AI screening tools solutions enabled by Intel technologies

9. What’s next

References

3

3

4

5

6

6

7

7

7

7

7

7

8

9

10

11

11

High-throughput screening (HTS) tools and clinical/lab screening workflows generate massive volumes of assay and quality-control (QC) data under strict turnaround-time constraints. Yet labs continue to face challenges in managing this data effectively, detecting subtle instrument drift, and rapidly identifying biologically meaningful hits. This solution blueprint introduces an Edge AI approach powered by heterogeneous compute (CPU + integrated GPU/NPU), enabling real-time, on-device analytics without reliance on the cloud. By intelligently partitioning workloads — leveraging CPUs for orchestration and control, integrated GPUs for parallel signal and image processing, and NPUs for efficient AI inference — the solution continuously monitors instrument health, performs automated QC (auto-QC), flags anomalies, and ranks samples by biological relevance. An integrated dashboard provides immediate visibility into system performance, top hits, and potential issues. The result is a scalable, low-latency architecture that transforms high-volume screening data into actionable insights, improves lab efficiency, reduces missed discoveries, and ensures data privacy while fully utilizing existing edge hardware infrastructure.

Edge AI architectures that combine CPUs with integrated accelerators such as iGPUs and NPUs are becoming foundational in modern life sciences devices, particularly for high-throughput screening devices like immunoassay analyzers, clinical chemistry systems, and molecular diagnostics. By enabling low-latency, on-device inference and efficient heterogeneous compute, these systems can process complex assay signals in real time while maintaining high throughput and performance-per-watt. Keeping data local enhances privacy and regulatory compliance, while eliminating cloud dependency improves reliability and uptime in clinical environments. At the same time, distributing workloads across CPU, iGPU, and NPU resources reduces cost and supports scalable deployment across instruments. This architecture also unlocks advanced edge capabilities such as real-time quality control (QC), anomaly detection, drift monitoring, and predictive maintenance — delivering the responsiveness and intelligence required for next-generation diagnostic and screening tools.

This solution blueprint is a technical guide for buyers at end-customer organizations, original design manufacturers (ODMs), system integrators, and device manufacturers. It details solution capabilities, specifications, and integration steps to support evaluation, informed adoption decisions, and a smooth path to deployment.

Revision Number	Description	Revision Date
1.0	Initial release	July 2026

1. Executive summary

High-throughput screening (HTS) and clinical/lab screening workflows generate large volumes of assay results and quality-control (QC) signals under tight turnaround-time expectations. This blueprint highlights an Edge AI solution for screening tools that shows how **Edge AI** running on **Intel® Core™ CPU + integrated Xe GPUs** can continuously monitor stability, flag anomalies in real time (auto-QC), prioritize likely “hits,” and surface operational insights — without sending sensitive data to the cloud.

2. Market analysis

The life sciences and diagnostics industry is undergoing rapid transformation driven by increasing test volumes, automation, and the growing complexity of assays across high-throughput screening (HTS) tools, immunoassay, clinical chemistry, and molecular diagnostics devices. Modern devices routinely process thousands of samples per hour, generating not only primary assay outputs but also extensive quality-control (QC), calibration, and system health data. While this data explosion creates opportunities for deeper insights, it also exposes a critical bottleneck.

Most laboratories lack the infrastructure and tools to efficiently process, analyze, and act on this data in real time. A key industry challenge is **data management at scale**. Traditional architectures rely heavily on centralized or cloud-based systems for downstream analytics, introducing latency, bandwidth constraints, and increased costs. This model is increasingly misaligned with clinical and laboratory requirements, where rapid turnaround times and continuous operation are essential. Additionally, the sensitive nature of patient and assay data imposes strict regulatory requirements, making data movement and external processing more complex and risk-prone.

Equally pressing is the challenge of **operational efficiency and instrument reliability**. Subtle issues such as reagent variability, sensor drift, or environmental fluctuations can degrade assay quality over time. These anomalies often go undetected until they impact results, leading to reruns, reduced throughput, and potential diagnostic inaccuracies. Existing QC processes are largely reactive and rule-based, limiting their ability to detect complex, multivariate patterns indicative of early system degradation. At the same time, identifying biologically meaningful “hits” from large screening datasets remains resource-intensive and time-consuming, slowing down research and clinical decision-making.

These challenges are unfolding within a rapidly expanding market opportunity for edge-enabled AI in healthcare and diagnostics. The **global healthcare Edge AI market** was valued at approximately **2.1 billion in 2025 and is projected to exceed \$10.4 billion by 2033, growing at 22% CAGR**, [1] reflecting strong demand for localized AI-driven analytics. In parallel, the **edge computing in healthcare market is expected to grow from \$8.2 billion in 2025 to \$47 billion by 2035 (19% CAGR)**, [2] driven by the need for real-time processing of high-volume medical data streams. More specifically, segments aligned with screening instruments — such as **Edge AI diagnostics platforms** — are projected to grow at **about 24% CAGR, reaching \$5.5 billion by 2035**, [3] while the underlying Edge AI accelerator market in healthcare [1] is expanding even faster (30% CAGR), underscoring the importance of heterogeneous compute (CPU, GPU, NPU) in enabling these workloads.

Despite this strong growth, **AI adoption within core screening and diagnostic workflows remains limited and uneven**, revealing a significant gap between potential and deployment. While a majority of healthcare organizations report some level of AI usage, adoption is largely concentrated in administrative functions or imaging-centric applications such as radiology. [4,5] Screening instruments remain largely dependent on rule-based logic and post-run analytics, with minimal integration of real-time, AI-driven decision-making.

This gap is particularly pronounced in high-throughput screening tool environments:

- Continuous generation of **high-volume, multimodal data streams** (signals, images, QC metrics).
- Reliance on **reactive QC workflows** rather than predictive or adaptive systems.
- Limited capability for **real-time anomaly detection, drift monitoring, or automated hit prioritization** at the instrument level.

At the same time, several barriers have slowed adoption, including integration complexity with legacy systems, regulatory and validation requirements, and lack of real-time processing infrastructure. However, advances in edge computing and heterogeneous architectures are now addressing these constraints, making it feasible to deploy AI directly on screening instruments without requiring cloud connectivity or major workflow disruption.

Market trends increasingly point toward **real-time, decentralized, and intelligent analytics embedded at the edge**. Vendors and laboratories are prioritizing solutions that improve throughput, enable proactive quality control, reduce total cost of ownership, and ensure data privacy through on-device processing. Edge AI powered by CPUs combined with integrated GPUs and NPUs is particularly well suited to meet these needs, enabling efficient workload partitioning and high-performance inference within existing hardware footprints.

3. Technical requirements

From a system architecture perspective, next-generation screening tool devices must be designed to handle **high-throughput, low-latency data processing** while improving **data management and operational efficiency** within constrained edge environments.

The key technical requirements are:

- **High-throughput data management:** Continuous ingestion, filtering, and processing of large-scale assay and QC data streams. Modern instruments generate continuous streams of assay signals, QC metrics, and telemetry that must be structured and managed without creating data overload for operators.
- **Real-time analytics:** Low-latency, on-device inference to meet turnaround-time requirements. To stay aligned with clinical and laboratory workflows, insights must be generated at the point of measurement — without waiting on centralized or cloud post-processing.
- **Operational efficiency:** Automated QC, drift detection, and reduced reruns to maximize instrument uptime. AI-driven monitoring can catch subtle degradation (calibration drift, reagent variability, signal noise) earlier than reactive, rule-based QC — reducing rework and protecting throughput.
- **Automated prioritization:** AI-driven ranking of hits and anomalies to reduce manual review burden. High data volumes create a manual triage bottleneck; automated ranking helps surface the most biologically meaningful hits and highest-risk anomalies for immediate review.

- **Edge data locality:** On-device processing for privacy, compliance, and reduced latency. Keeping sensitive assay and patient-adjacent data local simplifies regulatory alignment, reduces bandwidth dependency, and improves reliability in always-on clinical environments.
- **Heterogeneous compute utilization:** Optimized workload partitioning across CPU, iGPU, and NPU for efficiency. Many platforms already include integrated accelerators that are underutilized; splitting control and orchestration to the CPU and parallel inference to iGPU/NPU improves performance-per-watt and performance-per-dollar without new hardware.
- **Observable system performance:** Integrated telemetry and dashboards for monitoring throughput, latency, and resource utilization. Built-in observability (e.g., QC trends, anomaly tracking, hit scores, and CPU/iGPU utilization) ties architectural choices to measurable outcomes and supports validation during live execution.

In this solution blueprint, these key requirements are implemented through a streaming HTS-like pipeline that ingests assay outputs, QC metrics, and instrument telemetry. The system performs real-time feature extraction and executes two primary inference paths: a CPU-oriented auto-QC pipeline for anomaly detection and stability monitoring, and an iGPU-accelerated pipeline for hit scoring and prioritization. This demonstrates practical workload partitioning on commodity edge hardware while maintaining high throughput and responsiveness.



4. Target workflows (HST use cases)

Modern high-throughput screening tools (HST) and laboratory automation devices operate in environments where **speed, accuracy, and operational consistency directly impact clinical and research outcomes**. Edge AI enables a new class of real-time capabilities embedded directly within these systems, transforming them from passive data generators into intelligent, adaptive decision-support platforms.

To contextualize the impact, consider a routine clinical lab visit: a patient provides a blood sample for tests such as A1C, glucose levels, complete blood count (RBC and WBC), pregnancy screening, or disease markers. In parallel, pharmaceutical and biotech labs run high-throughput screening (HTS) assays to evaluate thousands of compounds for drug discovery. In both cases,

Use case	What it does	Why it matters
1. Real-time hit identification & ranking	Scores and ranks samples and compounds continuously (e.g., top 1% hits) during the run.	Speeds decisions in clinical interpretation and drug discovery by prioritizing the right results early.
2. Instrument health monitoring & auto-QC	Detects drift and subtle degradation using multivariate anomaly detection on telemetry + QC.	Reduces reruns and protects throughput by catching issues earlier than static rule-based thresholds.
3. Adaptive compute scaling	Shifts parallel workloads to iGPU/ NPU as demand rises while CPU handles control and orchestration.	Maintains predictable latency and higher throughput using existing heterogeneous compute.
4. On-instrument decision & QC dashboard	Displays hits, trends, anomalies, and health metrics locally with low latency.	Enables fast on-site action and uninterrupted operation without cloud dependency.

the core objective is the same — **deliver accurate results faster while maintaining confidence in data quality and system reliability**. Edge-AI-enabled screening systems directly address this need by accelerating analysis, improving QC rigor, and reducing time-to-insight at the point of measurement.

Within this context, representative Edge AI use cases include:

Impact of Edge-AI-enabled screening tools

Across both clinical diagnostics and high-throughput research environments, these capabilities converge to deliver:

- Faster turnaround from sample acquisition to result interpretation.
- Improved confidence through continuous, AI-enhanced QC.
- Reduced manual review and triage burden.
- Higher instrument utilization through adaptive compute scaling.
- Fully on-device intelligence without cloud dependency.

Ultimately, Edge AI transforms screening systems into **real-time analytical engines**, enabling clinicians and researchers to obtain reliable results faster — whether interpreting routine blood diagnostics or accelerating discovery in large-scale compound screening campaigns.

5. Reference architecture

5.1 System-level architecture:

Data pipeline architecture: The reference screening-tool design is a per-instrument, streaming pipeline that converts continuous, multi-modal instrument outputs (assay signals, QC metadata, calibration context, and telemetry) into real-time decisions at the edge. Incoming records are buffered to keep up with high-rate streams, aligned into traceable time windows, and transformed via feature engineering into model-ready tensors.

Inference is split across heterogeneous compute to meet latency and throughput targets: CPU-optimized auto-QC executes low-latency anomaly and drift checks, while the integrated GPU (and optionally an NPU, when available) runs parallel hit scoring and ranking. A decision engine fuses QC status and hit scores into alerts and prioritized candidate lists, and the same pipeline emits operational telemetry (throughput, latency, CPU/iGPU utilization) to support closed-loop monitoring and benchmarking.

Figure 1 illustrates this end-to-end streaming dataflow and the primary stage boundaries (ingestion and buffering, feature engineering, dual-path inference, fusion, and operator-facing telemetry).

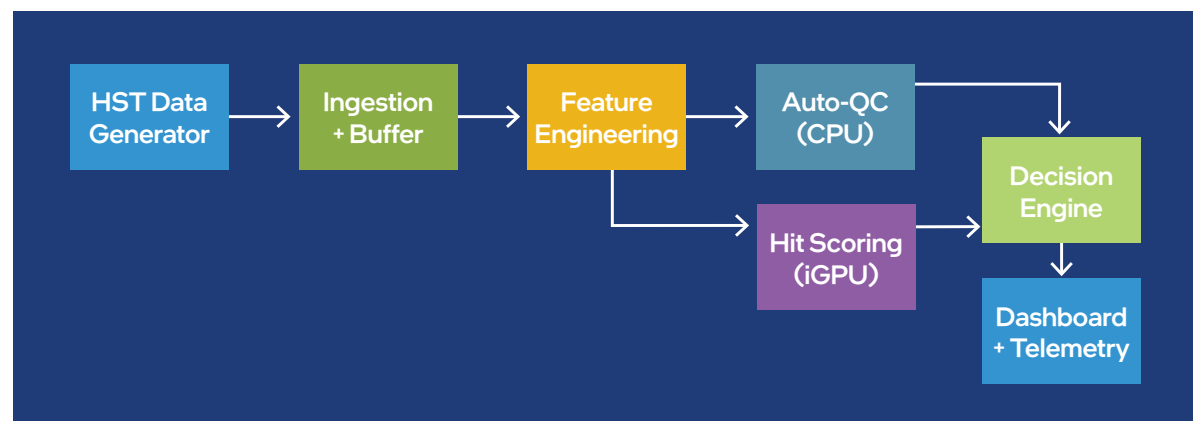


Figure 1: Screening tools conceptual data flow pipeline

From data flow pipeline to full system architecture

Figure 1 defines the logical flow of data and decisions through the Edge AI screening pipeline. Figure 2 expands that flow into implementation subsystems — multi-modal data sources, pipeline containerization per instrument, heterogeneous scheduling and runtime abstraction, decision and control-plane feedback loops, and real-time operational telemetry.

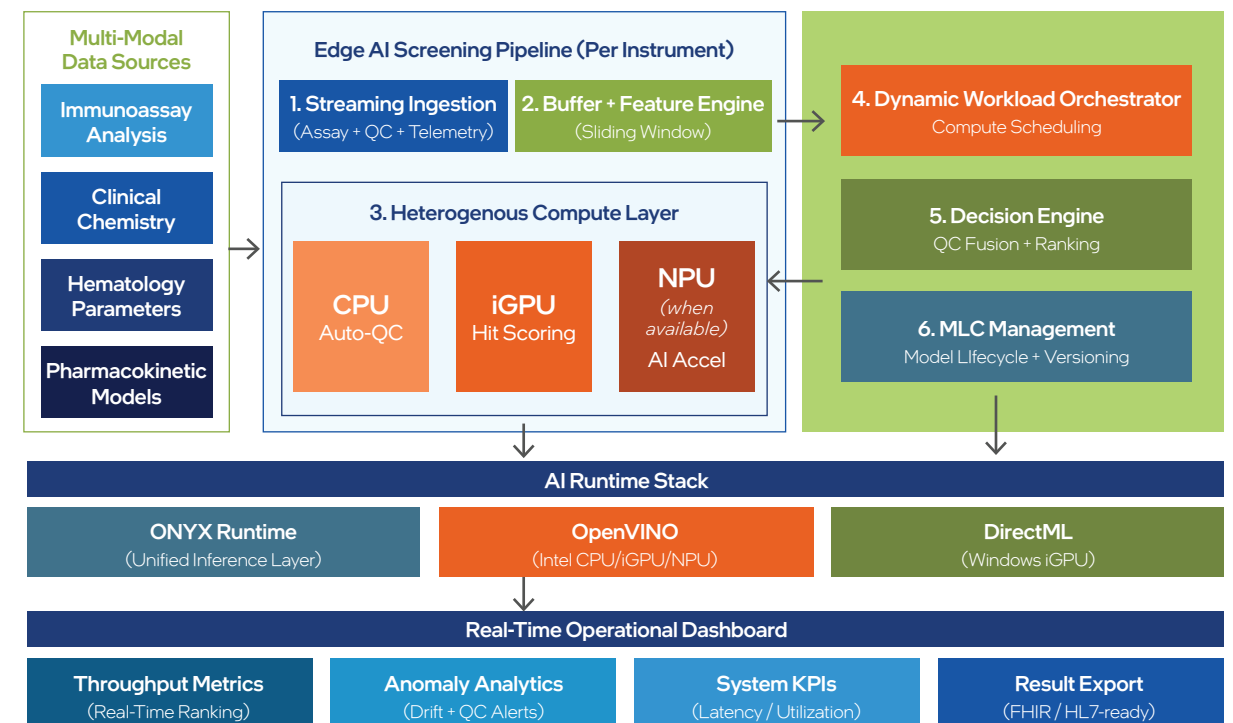


Figure 2: Heterogeneous compute architecture at full system level for screening tools

The next section decomposes these system-level blocks into concrete components, interfaces, and workload-placement choices that can be implemented and benchmarked on a representative Intel CPU + iGPU edge platform.

5.2 Component-level architecture

5.2.1 Reference system components

- **Data Generator (HTS Simulator)**
Produces synthetic assay streams, QC signals, and optional instrument telemetry at configurable throughput rates. Enables deterministic replay and stress testing.
- **Ingestion Layer**
Implements lightweight buffering (in-process queue or FastAPI endpoint) with backpressure control. Ensures continuous ingestion under burst conditions.
- **Feature Engineering Layer**
Computes rolling statistics, drift indicators, variance profiles, and model-ready tensors. Maintains per-channel sliding windows for temporal stability analysis.
- **Auto-QC Inference (CPU Path)**
Executes low-latency anomaly detection models optimized for CPU execution. Responsible for per-sample QC validation and early drift detection.
- **Hit Scoring Inference (iGPU Path)**
Executes batched MLP inference using ONNX Runtime with DirectML (or equivalent execution provider). Optimized for parallel compute on integrated GPU.
- **Decision Engine**
Aggregates outputs from both inference paths into structured artifacts:
 - Ranked hit list (e.g., top 1%).
 - Anomaly events with reason codes.
 - Aggregated KPIs (throughput, error rate, stability metrics).
- **Dashboard Layer (Streamlit UI)**
Provides real-time visualization, including assay trends, anomaly overlays, hit ranking views, and compute utilization telemetry.

5.2.2 Functional requirements

- Real-time ingestion of assay, QC, and optional instrument telemetry streams.
- Continuous feature computation and sliding-window analytics.
- Dual inference paths for:
 - Auto-QC anomaly detection (CPU).
 - Hit scoring and ranking (iGPU-accelerated).
- Real-time identification of top candidates (e.g., top 1% hits).
- Operator-facing dashboard with KPIs and interactive controls (batch size, mode switching).
- Workload scaling via batch-size adjustment to demonstrate iGPU utilization gains.
- Runtime observability including throughput, latency, and compute utilization.
- Exportable run summaries for benchmarking across configurations.

5.2.3 Non-functional requirements

- **Low latency:** Sub-second inference for QC and near-real-time ranking.
- **High throughput:** Sustain multi-thousand sample/hour workloads.
- **Deterministic behavior:** Repeatable synthetic runs with fixed seeds.
- **Portability:** Runs on commodity CPU + integrated GPU systems (Windows and Linux).
- **Data locality:** Fully edge-resident execution without cloud dependency.
- **Reliability:** Graceful handling of bursts, missing data, and backpressure conditions.
- **Resource efficiency:** Optimized use of heterogeneous compute (CPU/iGPU/NPU).
- **Observability:** Full system telemetry exposure for benchmarking and validation.

5.2.4 Inference execution and workload placement

- **Execution Providers (ONNX Runtime):**
 - CPUExecutionProvider → auto-QC pipeline.
 - DirectML/GPUExecutionProvider → hit scoring MLP.
- **Workload Partitioning Model:**
 - CPU: orchestration, feature extraction, QC inference.
 - iGPU: batched parallel inference (hit scoring).
 - Optional NPU: future acceleration for classification workloads.
- **Batching Strategy:**
Increasing batch size improves GPU utilization and throughput while introducing measurable latency trade-offs — used as a controlled scaling lever in the demo.
- **Concurrency Model:**
 - Producer thread: ingestion + feature computation.
 - Worker thread: inference execution.
 - UI thread: periodic state rendering (decoupled from inference loop).
- **Deterministic Controls:**
Fixed seeds and configurable thresholds ensure reproducible demo behavior across runs.

5.2.5 Observability and benchmarking

- The architecture exposes full-stack telemetry to validate performance and operational behavior:
- **Throughput Metrics**
 - Samples/sec.
 - Samples/hour steady-state.
 - **Latency Metrics**

- Feature computation latency.
 - Auto-QC inference latency.
 - Hit scoring latency.
 - End-to-end P50 and P95 latency.
 - **Compute Utilization**
 - CPU utilization (%).
 - iGPU or DirectML device utilization (%).
 - **Quality Metrics**
 - Anomaly detection rate.
 - Hit ranking stability across time windows.
 - Synthetic ground-truth precision and recall (demo mode).
 - **System Health**
 - Backpressure events.
 - Queue depth.
 - UI refresh lag.
- This reference architecture demonstrates how **edge-native heterogeneous compute systems** can transform screening tools into real-time intelligent devices. By combining streaming ingestion, feature-driven analytics, and dual-path AI inference across CPU and integrated GPU, the system achieves:
- Real-time QC and anomaly detection (auto-QC).
 - Continuous hit ranking for biological relevance.
 - Efficient utilization of existing edge compute resources.
 - Full data locality for privacy and compliance.
 - Observable, benchmarkable system performance.
- This architecture provides a blueprint for evolving next-generation screening tools from static analyzers into **adaptive, AI-driven diagnostic systems operating entirely at the edge**.

6. Reference hardware configuration — screening tools

For the solution blueprint, the target deployment is a single **edge desktop computer** (or future on-instrument compute) with a modern Intel desktop CPU and integrated GPU. This mirrors realistic lab environments where a discrete GPU may not be available or desirable. The blueprint emphasizes: **keeping inference local, using the iGPU for compute-intensive hit scoring, and reserving the CPU for orchestration, auto-QC, and I/O.**

- **CPU/iGPU:** Intel Core Series 2, Intel® Core™ Ultra Series 3, or newer (or equivalent Intel iGPU-capable platform).
- **Memory:** 16–32 GB RAM.
- **Storage:** SSD recommended for fast environment setup and model loading.
- **OS:** Windows or Linux.
- **Drivers and iGPU enablement:** Current Intel graphics driver (required for iGPU acceleration with ONNX Runtime DirectML on Windows). Confirm the iGPU is enabled in BIOS or UEFI and visible in OS device inventory.
- **Discrete accelerators:** No discrete GPU required; the blueprint is designed for offline or local execution.
- **Power and thermals:** Use AC power and set OS power mode to “Best performance” (or equivalent).
- **Memory sizing guidance:** 16 GB is sufficient for the baseline demo; prefer 32 GB if you expect larger batch sizes, additional telemetry channels, longer rolling windows, or multiple concurrent runs. Ensure adequate headroom to prevent paging during steady-state throughput tests.
- **Storage capacity and I/O:** Keep at least 5–10 GB free for Python environments, model files, logs, and exported run summaries. SSD is recommended to reduce environment setup time and avoid slow model load and startup.

- **Networking (optional):** If there is a stream of data from another device or a need to connect to an instrument emulator over LAN, ensure the network is trusted and stable (and document required ports and addresses).
- **Hardware validation checks (recommended):** Confirm the iGPU is recognized (e.g., Windows Task Manager → Performance → GPU) and that utilization changes when batch size increases. If iGPU is not visible or stays at 0%, verify BIOS iGPU enablement and update the graphics driver.

Workload distribution in reference solution:

Hardware	Workload	Description
CPU	Auto-QC	Continuous quality control checks and anomaly detection on incoming assay data.
	Baseline inference	Running the initial scoring or pre-processing steps for samples before AI hit scoring.
iGPU	Hit scoring (deep MLP)	Predicting AI hit scores for prioritizing biologically meaningful samples.
	Neural network inference	Any additional deep learning inference, e.g., predicting assay trends, ranking, or enhanced analytics.
NPU (when available)	Inference	Offloading any additional or new inference workloads to NPUs.

7. Data handling, privacy, and governance

For high-throughput screening (HTS) tools, the safest and most scalable default is to keep assay, QC, and telemetry data local and minimize what leaves the instrument or lab network. On Intel edge platforms, this is enabled today by running analytics and inference on-device (CPU + integrated GPU/NPU as available), using OS and platform security capabilities for access control and least-privilege execution, and protecting stored artifacts with encryption at rest (commonly OS full-disk or file encryption accelerated by Intel® Advanced Encryption Standard (Intel® AES) instructions). Governance should treat logs and exports as data products: prefer aggregated

KPIs over per-sample records, avoid raw-record logging, and apply explicit “demo/synthetic” vs. “internal/real-data” modes that tighten defaults (export off, stricter logging, enforced encryption). For traceability, keep configuration and model artifacts version-pinned (hashes and versions recorded per run) so results are reproducible and changes are auditable — critical in HTS environments where throughput, uptime, and compliance requirements make uncontrolled data movement and ad-hoc sharing a risk.



8. Potential Edge AI screening tools solutions enabled by Intel technologies

Potential Solution	Intel Hardware & Software Stack	Details (Architecture + Value Enablement)
Real-Time Auto-QC for Immunoassay & Clinical Chemistry Analyzers	Intel® Core™ processors + Intel® Xe iGPU + Intel® Distribution of OpenVINO™ Toolkit	▪ Enables low-latency anomaly detection directly on instrument streams (assay signals, QC metrics). ▪ CPU handles control + feature extraction, iGPU accelerates parallel inference. ▪ OpenVINO optimizes model execution for deterministic sub-second QC decisions, reducing reruns and improving instrument uptime.
High-Throughput Hit Scoring for Drug Discovery (HTS Platforms)	Intel® Core™ Ultra (CPU + integrated NPU + iGPU) + OpenVINO + ONNX Runtime	▪ Supports large-batch inference for compound ranking and biological signal prioritization. ▪ iGPU/NPU accelerates MLP/CNN inference workloads, while CPU manages orchestration. ▪ Enables real-time “top 1% hit” identification at instrument edge without cloud dependency.
Molecular Diagnostics Signal Interpretation & Classification	Intel® Core+ Intel Xe iGPU + Intel® OpenVINO Runtime	▪ Provides high-accuracy classification of genomic and amplification signals and molecular assay outputs. ▪ Suitable for multi-instrument aggregation at lab edge. ▪ Improves diagnostic consistency and reduces manual review burden in high-complexity assays.
Edge-Based Drift Detection & Predictive Maintenance for Lab Instruments	Intel Core processors + Intel® oneAPI + OpenVINO	▪ Runs continuous time-series anomaly detection on instrument telemetry (temperature, motor current, reagent variability). ▪ Enables early detection of calibration drift and mechanical degradation, improving uptime and reducing maintenance costs.
LIS/ LIMS Integrated Edge Data Management for Screening Workflows	Intel® Edge Platform (Core + Xeon) + Intel® Smart Edge + FHIR/HL7-enabled middleware	▪ Provides structured ingestion of assay + QC outputs into LIS/LIMS systems. ▪ Ensures standardized data contracts across instruments and supports compliance-ready traceability (audit logs, PHI locality). ▪ Reduces fragmentation between instruments and lab IT systems.
Fleet-Wide Multi-Instrument Screening Orchestration (Lab Network Scale-Out)	Intel® Xeon® Servers + Intel® Ethernet + Kubernetes on Intel Architecture	▪ Enables centralized orchestration across multiple analyzers (immunoassay, chemistry, HTS). ▪ Aggregates KPIs (throughput, QC anomalies, latency) across lab fleet. ▪ Supports workload balancing and model distribution across distributed lab environments.
Hybrid Edge-Cloud Model Training & Model Update Distribution	Intel® OpenVINO Model Server + Intel® Xeon® Scalable Processors	▪ Supports federated learning and periodic model retraining using de-identified datasets. ▪ Updated models are optimized for edge deployment via OpenVINO and distributed back to instruments, enabling continuous performance improvement without disrupting clinical operations.

9. What's next

What can you do next:

- Evaluate next-gen **Intel Core Series (and Intel Core Ultra)** platforms to modernize on-instrument or near-instrument compute for higher-throughput data handling and improved operational efficiency.
- Enable and benchmark AI workloads on your **existing Intel® Xe integrated GPU (iGPU)** (and NPU where available) to accelerate hit scoring, auto-QC, and telemetry — **without adding new hardware**.
- Contact the Intel team to see a **working demo** of the end-to-end solution described in this blueprint, including live dashboard functionality.
- Engage Intel for a focused evaluation or **proof of concept** to map the reference pipeline to your instrument data flows and define target KPIs (throughput, latency, QC sensitivity, and utilization).

References

Citation	Description
1	GrandView Horizon: Healthcare – Edge AI market statistics, 2025-2033
2	Precedence Research: Edge Computing in Healthcare Market
3	EIN Presswire: Edge AI Diagnostics Platform Market
4	PubMed Central: JAMIA: Adoption of AI in healthcare
5	Menlo Ventures: 2025: The State of AI in Healthcare





Notices & Disclaimers

You may not use or facilitate the use of this document in connection with any infringement or other legal analysis concerning Intel products described herein. You agree to grant Intel a non-exclusive, royalty-free license to any patent claim thereafter drafted which includes subject matter disclosed herein.

No license (express or implied, by estoppel or otherwise) to any intellectual property rights is granted by this document.

All information provided here is subject to change without notice. Contact your Intel representative to obtain the latest Intel product specifications and roadmaps.

The products described may contain design defects or errors known as errata which may cause the product to deviate from published specifications. Current characterized errata are available on request. Copies of documents which have an order number and are referenced in this document may be obtained by calling 1-800-548-4725 or visiting the [Intel Resource and Documentation Center](#).

Performance varies by use, configuration, and other factors. Learn more at www.Intel.com/PerformanceIndex. Differences in hardware, software, or configuration will affect actual performance. Your results may vary.

Intel technologies may require enabled hardware, software, or service activation. Learn more at intel.com, or from the OEM or retailer. Performance results are based on testing as of dates reflected in the configurations and may not reflect all publicly available updates. See configuration disclosure for details.

Intel is committed to respecting human rights and avoiding complicity in human rights abuses. See Intel's Global Human Rights Principles. Intel's products and software are intended only to be used in applications that do not cause or contribute to a violation of an internationally recognized human right. Intel technologies' features and benefits depend on system configuration and may require enabled hardware, software, or service activation. Performance varies depending on system configuration. No product or component can be absolutely secure. Check with your system manufacturer or retailer or learn more at intel.com.

© Intel Corporation. Intel, the Intel logo, and other Intel marks are trademarks of Intel Corporation or its subsidiaries. Other names and brands may be claimed as the property of others.

0726/EC/MESH/PDF 367776-001US