

Hardware and performance

CPU, RAM, SSD, network; measured Analyze times; PC and file-server sizing.

This guide is for IT and purchasing. It explains what the Analyze pipeline actually does, which hardware limits each step, typical times on a **20-logical-processor** NVMe laptop, and how to size workstations, laptops, and file servers.

Quanto is a **local desktop** process. There is no Quanto application server. A “server” in this document is a **file server** (or instrument PC) that holds raw data, methods, and finished `.qtb` files.

The **reference times** in this guide are from a measured Analyze on the 20-thread laptop (Quanto 1.0.0.8, 2026-09-05). Other batch sizes are scaled from that run. Logbook scopes to copy when you bench a new PC: `Load from data provider`, `Integrate`, `Identify`, and `Calibrate`, `quantitate` and `outliers`.

Reference laptop (20 threads)

Quanto sets `MaxDegreeOfParallelism = logical processors - 2` (default **on**). On a 20-thread PC that is **18 workers**. Two threads stay free for the UI.

Resource	What “20-core laptop” usually means	Why it matters
CPU	20 logical processors (often 10–14 physical cores with hyper-threading, or 20 threads on a recent Intel / AMD mobile chip)	Extract and integrate scale across samples. Calibrate / quantitate do not.
RAM	32 GB is the practical floor (this 104,544-signal batch needs it); 64 GB if you hold several such batches or much wider windows	The whole batch’s chromatograms stay in memory.
System disk	NVMe SSD, about 3–7 GB/s sequential, 200k–500k random 4K IOPS	Extract is random reads of vendor files. Save/open <code>.qtb</code> is sequential zip I/O.
Network	1 Gbit Ethernet $\approx$ 110 MB/s best case; Wi-Fi is lower and less stable	A MassHunter <code>.d</code> folder is thousands of small files. SMB latency dominates extract.

If Task Manager shows 20 logical processors, you match this column. Physical “20-core” desktop chips are faster per thread than thin laptops; the **scaling** (what gets faster when you add cores) stays the same.

What Analyze does (and what it waits on)

Analyze (full batch) runs in this order:

Step	Parallel?	Bound by	What it does
Get signals (extract)	Yes — one sample per worker (up to 18). Agilent also prefetches indexes (capped around 8 opens).	Disk / network , then CPU	Opens each <code>.d</code> / <code>.wiff</code> , reads TIC + every method channel (quant + qualifiers), optional smooth.
Integrate	Yes — samples, and compounds inside a sample. Per-thread integrator instance.	CPU	Peak detect + baseline on every chromatogram.
Identify	Yes (sample loop)	CPU	Pick peaks, qualifiers, ISTD links, responses.
Calibrate	No (per quant / bracket)	CPU (single thread, curve engine)	Fit each curve. Small compared with extract.
Quantitate	No (walk every result)	CPU (single thread, cheap math)	Inverse curve + dilution + accuracy. Almost free.
Save <code>.qtb</code>	Chromatogram buffers written in parallel; zip is one file	Disk	Compress tables + all chromatograms + integrations + optional report <code>xlsx</code> .
Open <code>.qtb</code>	Read zip, inflate chromatograms	Disk / network	Reload results without touching raw files.
Report export	Up to 4 STA threads	CPU + disk	One Excel/PDF pair per sample.

Re-Analyze after you already have signals skips extract and is usually **integrate + identify + calibrate** only.

Measured reference (20-thread laptop, local disk)

Batch: Agilent MassHunter LC-MS pesticides demo (`D:\MassHunter\Data\Demo\LCMS_Pesticides\`). Status bar after Analyze (Quanto 1.0.0.8): **99 samples, 461 analytes, 104,544 signals, 16,837 peaks**, 1 method. Saved batch `2026-09-05.qtb` = **169,324 KB (~165 MB)**. Indexed `.d` data on a local drive. Parallel processing on (18 workers).

That file is a zip (Explorer may show a 7-Zip icon). On disk it is about **1.6 KB per signal** or **1.7 MB per sample**. RAM after open is several times larger because chromatograms are inflated to `double` arrays.

Use the status-bar **Signals** count for hardware sizing. **Peaks** (16,837) are detected peaks, not chromatograms. This run is ~1,056 signals/sample and ~2.3 channels per analyte per sample.

Full Analyze (ProcessBatch = 6.72 s)

Step (Logbook message)	Time	Share	Rate
Load from data provider (extract)	3.23 s	48%	33 ms/sample · 32,000 signals/s
Integrate	1.97 s	29%	20 ms/sample · 53,000 signals/s
Identify	0.77 s	11%	8 ms/sample
Reapply manual overrides	0.10 s	1%	
Calibrate, quantitate and outliers	0.46 s	7%	
Choose reported sample	0.12 s	2%	
UI grid commit (Results)	0.04 s		
ProcessBatch total	6.72 s	100%	68 ms/sample · 16,000 signals/s

Calibrate / quantitate / outliers inside that 0.46 s:

Substep	Time
Calibrate	66 ms
Recompute auto quant selection	51 ms
Quantitate	134 ms
Outliers (99 samples)	138 ms

Identify is mostly **Assign method items** (457 ms). Correlate, qualify, peak pick, co-elution, and ISTD link are tens of milliseconds each.

Open saved batch and start

Operation	Time	Notes
Start Quanto (paths + settings + add-ins + main docking)	~8 s	Add-ins 0.39 s; docking LoadSettings 1.7 s
Open .qtb (OpenOrCreateBatch)	6.2 s	Inflate 165 MB / 104,544 chromatograms from the zip (~27 MB/s)
Calibrate + quantitate after open	1.1 s	Calibrate 638 ms, quantitate 206 ms, outliers 184 ms

Opening this batch is about as expensive as Analyze. Review PCs need a fast local disk too, not only extract PCs.

Scale from this run (local SSD, same method style)

Linear in **signal count** for a first IT estimate. Extract grows faster than linear on a slow network; integrate stays close to linear on CPU.

Workload	Signals	Extract	Integrate + identify	Cal + quant	Full Analyze
Measured — 99 × 461, 16,837 peaks	104,544	3.2 s	2.7 s	0.46 s	6.7 s
Same method, ~24 samples	~25k	~0.8 s	~0.7 s	< 0.2 s	~1.6 s
Same method, ~50 samples	~53k	~1.6 s	~1.4 s	~0.2 s	~3.4 s
Same method, ~200 samples	~211k	~6.5 s	~5.5 s	~0.9 s	~14 s
2× this channel count	~209k	~6.5 s	~5.5 s	~0.9 s	~14 s
Same 104,544 signals from a 1 Gbit share	104,544	15–50 s (estimate)	2.7 s	0.46 s	20–55 s

SCIEX WIFF is often faster to extract over the network than Agilent .d (fewer large files). Full-scan EIC extract is slower than this MRM demo.

First-time Agilent index conversion (TDA) is extra and writes MSTree2.bin into each .d. Plan tens of seconds to a few minutes per sample the first time. This measured run was already indexed.

How each resource changes the number

CPU

- **Helps:** extract (after the file is in cache), integrate, identify, report pictures, UI charting.
- **Does not help much:** calibrate, quantitate, zip save, opening a huge .qtb.
- **Diminishing returns above ~16–20 threads** for extract: prefetch already caps Agilent opens, and disk queue depth saturates.
- Prefer **high per-core speed** (recent laptop/desktop, not a many-core server chip with low clocks) for integrate and the UI.
- Leave **2+ cores free** (the default processors - 2). Setting parallelism to “all cores” makes the window stutter.

Purchase: **16–20 threads** is the sweet spot for an analyst PC. **8 threads** is usable for small batches. **32+ threads** is only worth it if you also have local NVMe and large residue methods.

RAM

Chromatograms stay in RAM as double arrays (time + response, original and often a modified copy):

Points per chromatogram	Bytes per signal (orig + modified)	105k (measured)	210k	400k
500 (narrow MRM window)	~16 KB	1.7 GB	3.4 GB	6.4 GB
2 000 (typical)	~64 KB	6.7 GB	13 GB	26 GB
8 000 (wide window / scan)	~256 KB	27 GB	54 GB	100 GB

Add **1–3 GB** for the WPF UI, grids, docking, and report designer, plus TIC traces. The measured 104,544-signal batch is in the **~8–12 GB** working-set band at typical MRM point counts.

RAM	Comfortable batch
16 GB	Tight for this measured batch (105k signals). Possible if chromatograms are short and little else is open.
32 GB	Standard analyst laptop. Proven for 99 samples / 461 analytes / 104,544 signals .
64 GB	Several batches, Key Review, 4K + report designer, or ~200k+ signals.
128 GB	Multi-batch review or very wide scan windows.

Windows will use the page file if you undersize RAM. Analyze then becomes **disk-bound** and can look “stuck.” That is not a Quanto bug.

SSD vs HDD

Storage	Extract (local)	Save / open .qtb	Verdict
NVMe (Gen3/4), 3–7 GB/s	Baseline in the table above	Measured 165 MB .qtb opened in 6.2 s	Required for the system disk and for active studies
SATA SSD	~1.5–2× slower extract (random IOPS)	Fine	Acceptable for a second data disk
7200 rpm HDD	5–20× slower extract	.qtb save can take minutes	Archive only
USB stick / SD	Do not use	Do not use	

Random **IOPS** matter more than sequential GB/s for Agilent .d folders. A cheap SATA SSD still beats a fast HDD.

Put **Windows + Quanto + the active study** on NVMe. Cold archive can be HDD or NAS.

Network

Extract reads **many small files** (MassHunter AcqData , ChemStation). SMB latency per file adds up.

Link	Realistic extract vs local NVMe	When to use it
Local NVMe	1×	Daily Analyze
10 Gbit Ethernet, SSD/NVMe NAS, SMB3	1.2–2×	Shared study disk if the NAS is fast and antivirus is excluded
1 Gbit Ethernet	5–15× slower extract	OK for opening a finished .qtb or loading a .qtm . Poor for first Analyze of a large .d sequence
Wi-Fi 6	Unstable; often worse than 1 Gbit	Demo only. Copy the sequence local first
VPN / WAN	Often 30–100×	Do not extract. Copy data, or Analyze on a PC next to the files

Good on a share: .qtm methods, .lic licenses, finished .qtb for review, QuantoReports .

Bad on a slow share: first Get signals from Agilent .d trees, TDA conversion (it writes back into the sample folder).

Practical rule: **Analyze next to the data.** Copy the sequence to the laptop NVMe, or sit at the instrument PC. After save, copy the .qtb to the LIMS share.

Antivirus on the NAS or the endpoint that scans every .bin in AcqData can be as slow as a 1 Gbit link. See [exclusions](#).

GPU and display

Quanto does not use the GPU for extract, integrate, or quantitate. An **integrated GPU** is enough. A discrete GPU only helps if you use many 4K chromatogram plots or Remote Desktop with GPU encoding.

Two 4K monitors increase RAM a little (WPF visuals) but do not change Analyze time.

Batch size (what to tell the lab)

Size the PC by the status-bar **Signals** count, not Peaks and not samples × analytes . This 461-analyte batch is **104,544 signals** and 16,837 peaks.

Class	Signals	RAM	Disk for .qtb	20-thread local Analyze
Small	< 25k	16 GB	~40 MB	~2 s
Standard	25k–80k	32 GB	40–130 MB	3–6 s
Measured pesticides	104,544	32 GB	165 MB (2026-09-05.qtb)	6.7 s
Large	80k–200k	32–64 GB	130–330 MB	6–15 s
Extreme	> 200k	64–128 GB	330 MB–1 GB	15 s+; keep data local

Scaled .qtb sizes use **1.6 KB/signal** from the measured file. A report workbook inside the zip adds extra megabytes.

Key Review (several AcqKeys / QntKeys / dilutions) multiplies **rows in the UI**, not necessarily extract time, if the extra keys already have signals. Memory and grid refresh grow with visible cells.

Copying this **165 MB** .qtb on 1 Gbit takes a few seconds. Opening it still needs ~6 s locally to inflate. A 1 GB batch on a slow share can take longer to open than Analyze on local NVMe. Review PCs should copy large batches local or use 10 Gbit.

Workstation and laptop purchase

Quanto needs **Windows 10 or 11 64-bit**, .NET 10 Desktop Runtime (x64), and a local SSD.

Role	CPU	RAM	Storage	Network	Notes
Analyst laptop (matches the reference)	16–20 threads, recent generation	32 GB (64 GB if you routinely exceed ~200k signals)	1 TB NVMe (OS + active studies)	1 Gbit or Wi-Fi 6; dock to Ethernet for copy	Measured: 99 / 461 / 104,544 signals / 16,837 peaks; 165 MB .qtb ; 6.7 s Analyze, 6.2 s open.
Review / senior workstation	16–24 threads, high clocks	64 GB	2 TB NVMe	1–10 Gbit	Several batches open, report designer, Key Review.
Heavy multi-residue	20+ threads	64–128 GB	2 TB+ NVMe	10 Gbit to the data disk	Do not extract from Wi-Fi.
Instrument PC (if they Analyze there)	Vendor-min or better	32 GB	Fast local disk for the sequence	Same LAN as the lab	Best place to run TDA / first extract.
File server / NAS	Few cores are fine	32 GB+ cache	SSD/NVMe pool for hot studies; HDD for archive	10 Gbit recommended	SMB3. Exclude real-time AV on .d / .qtb . Not a Quanto host.
VDI / Citrix	8+ vCPU dedicated	32 GB session	Local or cache the study on the session host	Low latency to the host	GPU optional for UI. Never extract a .d tree from the user's home NAS over WAN.

Do not buy: 8 GB RAM PCs, HDD-only laptops, or a “Quanto server” VM with no GPU and data on another continent.

Nice but unused: ECC RAM, dual CPU, compute GPU, Linux (Quanto is Windows x64 only).

Server and lab network

1. **No application server.** Do not install Quanto on a headless VM expecting batch jobs. It is an interactive WPF app.

2. **File server** holds sequences, `.qtm`, `.qtb`, licenses. Prefer 10 Gbit and SSD for the live study volume.
3. **License share** (`*.lic`) can be a small, highly available folder. Analysts need read; one admin needs write. See [installation folders](#).
4. **One writer per** `.qtb`. Do not let two PCs save the same batch file at once. Review copies are fine.
5. **Instrument** → **process** → **archive**: acquire on the instrument disk → copy to analyst NVMe (or Analyze on the instrument PC) → save `.qtb` → copy `.qtb` + reports to LIMS / NAS → archive raw data per SOP.
6. **Domain / roaming**: roam only `%LOCALAPPDATA%\Quanto` settings if you must; never roam `BatchReports` or studies.

Checklist for a new PC

- ☐ Windows 10/11 x64, current .NET 10 Desktop Runtime
- ☐ 32 GB RAM minimum (64 GB if status-bar Signals regularly exceeds ~200k)
- ☐ NVMe system disk; active studies on that disk or another SSD
- ☐ 1 Gbit Ethernet at the bench (do not rely on Wi-Fi for extract)
- ☐ Antivirus exclusions for `C:\Program Files\Quanto` and study data
- ☐ Write access to Agilent `.d` folders if conversion may run here
- ☐ MassHunter Quant / `TDAConverterConsole.exe` only if instruments do not already write `MSTree2.bin`
- ☐ License folder reachable (Public Documents or a share)
- ☐ Confirm Task Manager → 16+ logical processors if you promised the times in this guide

Related

- [Installation and folders](#)
- Analyst workflow: [Set up an analysis](#)