



THE GTMIFY OPERATING SYSTEM

The Parallel Operator

How one operator runs a software company, its operations and a household with 24 to 30 parallel Claude Code sessions, a fleet of agents and a guarded toolchain. Written down in full.

BEFORE YOU START

What this is

This is the complete operating system I use to run GTMify: the software company, its operations, and, with the same pattern, my household. It is written down in full, and nothing is held back for a sales call.

On a normal day, 24 to 30 Claude Code sessions run at the same time on my machine. Each one owns a surface of the business. Builder agents take tickets to pull requests, reviewer agents on a different model read those pull requests adversarially, a merge train lands the work, and a release goes to production behind explicit preconditions. Outside the product, a control plane of more than 45 agents handles meeting prep, account research, deal reviews and campaign work. I supervise, decide, and keep the judgment calls.

The document runs in the order you would build it. Part 1 argues why this is a systems problem. Parts 2 through 5 describe the machine: where the capacity goes, how the lanes are organized, how a backlog becomes shipped code, and the toolchain underneath. Part 6 covers the guardrails and the failures that produced them. Parts 7 and 8 widen the frame to company operations and to the household. Part 9 is a four-week adoption path you can start on Monday.

A note on what happens next

Most readers will change nothing after reading this, and that is a reasonable outcome. The system took months of iteration, and every rule in it traces back to a failure that cost me real time. If you read it, take two habits from Part 9 and stop there, you will still get more out of the plan you already pay for.

If you want the result without the months, the last page says how to get it. That is the only offer in the document.

PART 1 · THE THESIS

The constraint is the operator

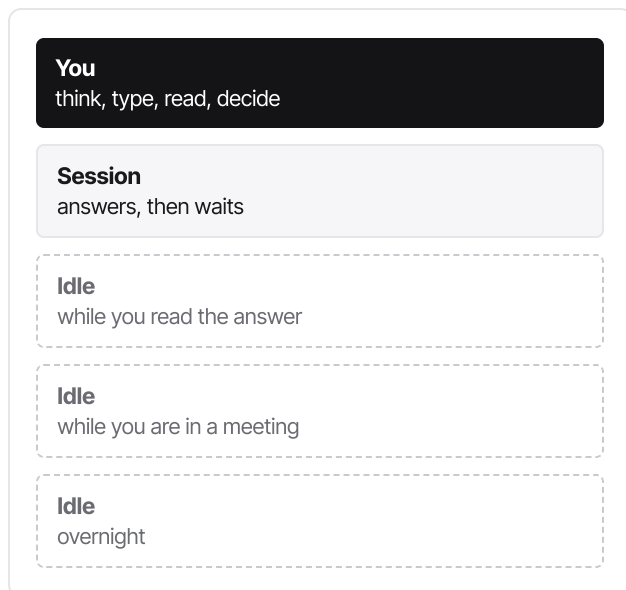
Most go-to-market strategies fail from the absence of a system, and the effort behind them is rarely the problem. AI adoption fails the same way.

I have spent my career engineering revenue engines, and the pattern is consistent. A team works hard, buys good tools, and still misses the number, because the tools, the people and the process were never connected into one workflow. Each part is competent and the whole is slow.

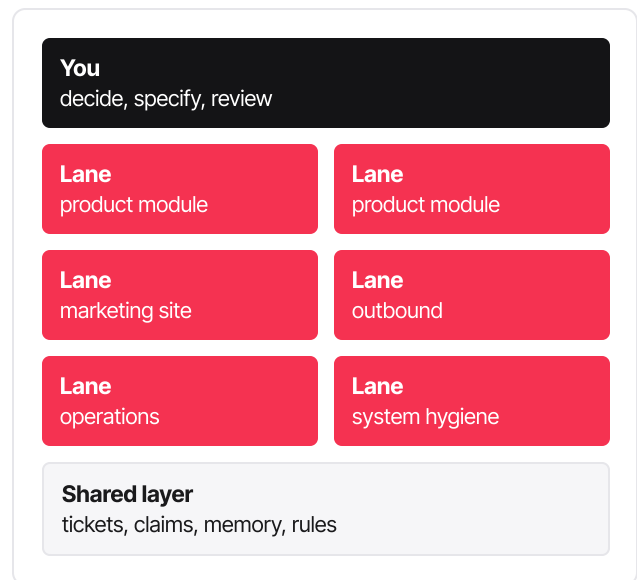
The way most people use AI reproduces that failure at the scale of one person. They open a chat, ask a question, read the answer, think, and ask the next one. The model is fast and capable, and it spends most of its time waiting for a human to finish reading. Upgrading to a bigger plan changes nothing about that loop, because the plan was never the bottleneck. **The operator's attention is the bottleneck**, and every hour of work still passes through it in sequence.

The whole of this document rests on one move: **remove yourself as the serial bottleneck by running work in parallel lanes, each owning one surface, while you supervise the set**. You stop being the person who does each task and become the person who decides what gets done, specifies it well enough to delegate, and checks what comes back.

SERIAL: ONE OPERATOR, ONE SESSION



PARALLEL: ONE OPERATOR, MANY LANES



Why this is a systems problem and not a prompting problem

Prompting advice improves one exchange. A system improves every exchange, including the ones you are not present for. The questions that decide how much work you get out of AI are structural: how work is divided so sessions do not collide, how each session knows the rules without being told again, how a finished piece of work

is checked by something other than the thing that wrote it, and how the human is asked for a decision without being buried in status.

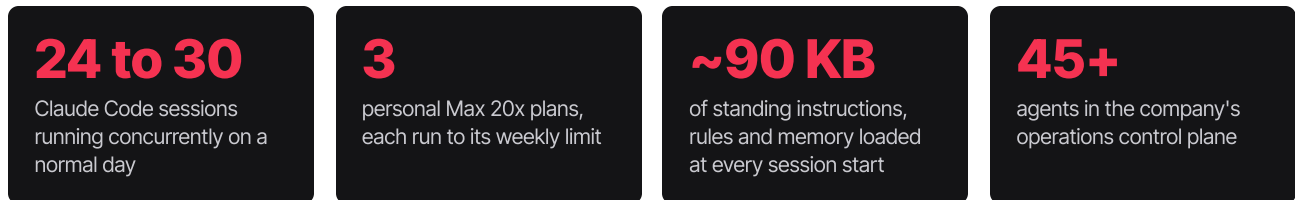
Each of those questions has an answer in the pages that follow. None of the answers is clever. What makes them work is that they are written down once, in files every session reads, and enforced by configuration where possible so they hold even when a session forgets.

Write the system down once, in files an agent can read. Keep the judgment for yourself.

PART 2 · WHERE THE CAPACITY GOES

What three plans a week actually buy

People who hear how I work usually ask the same question first: what could possibly consume three of the largest personal Claude plans every week? The answer is parallelism, and the arithmetic is simple.



A single operator working one session at a time rarely hits a limit, because their own reading and typing pace caps the work. My setup removes that cap. When one session is waiting on a build, dozens of others are still reading code, writing tickets, reviewing pull requests, researching accounts, or drafting outbound. Usage scales with **sessions multiplied by turns multiplied by context size**, and I push all three up on purpose.

The three multipliers

Multiplier	What drives it	Why the cost is worth paying
Sessions	One terminal per product module, plus lanes for the marketing site, the company's own outbound, a second operating company, and a hygiene lane that maintains the system itself.	Throughput. Work that used to take a small team a week lands in a day.
Turns	Sessions run long autonomous loops: plan a batch of tickets, dispatch builders, review, merge, verify, and repeat until the backlog is empty.	The loops run overnight and while I am in meetings, which is time a serial operator cannot use at all.
Context	Every session boots with the same standing instructions, rules and memory index, about 90 KB measured in September 2026. Every subagent loads its own context on top.	Consistency. Thirty sessions follow the same rules without anyone repeating them.

Spending deliberately

Pushing all three multipliers up only pays if each one is spent on purpose. Three disciplines keep it that way.

Standing context has a budget. The instructions every session loads at boot cost tokens on every turn of every session, so a paragraph added there is a paragraph thirty sessions pay for all day. A hook measures that layer at every session start and warns when it grows past its line. The rule that goes with it is one in, one out: a session

that adds a standing rule retires another in the same change, usually by moving evidence and history into a reference file that loads only when needed.

Delegation passes references, not contents. When a session hands work to a subagent, it passes file paths and acceptance criteria, and the subagent writes its summary to a file. Pasting file contents into a delegation doubles the cost of the work and saves nothing.

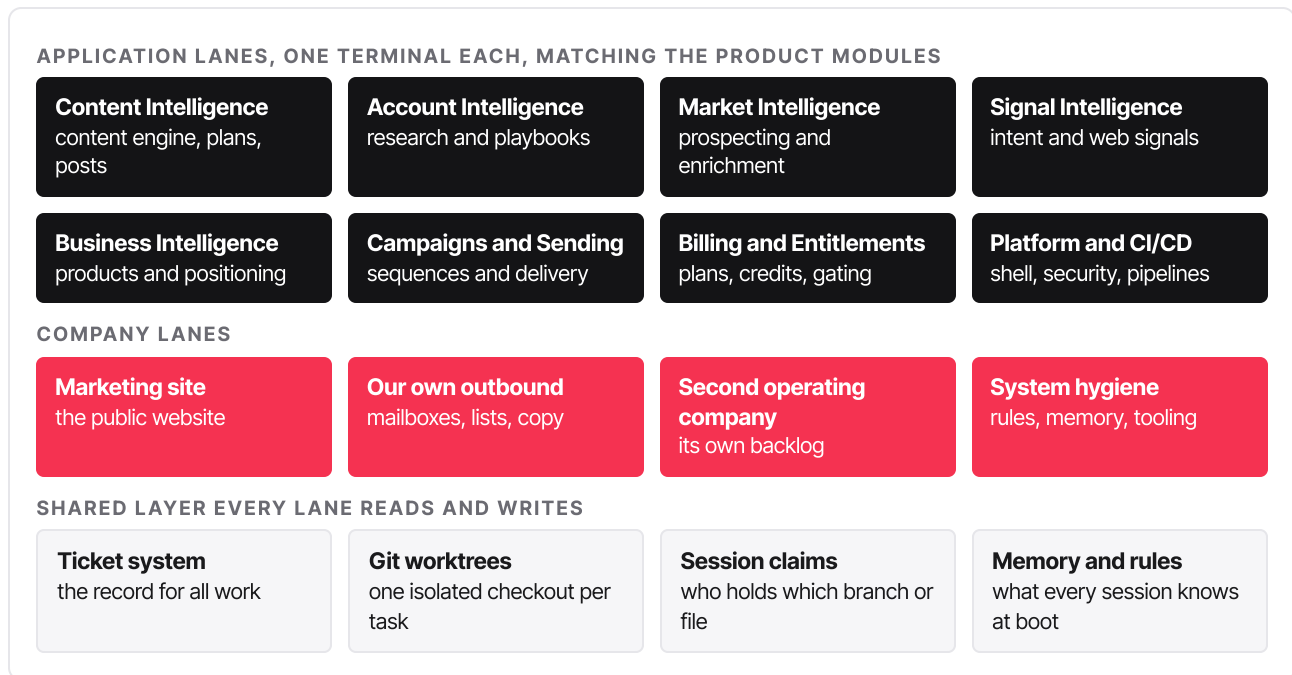
The model matches the job. The most capable model handles design, debugging, and anything touching money or permissions. Lighter models handle mechanical, fully specified edits, but only where a mistake would fail loudly. Where a wrong answer would look plausible, the better model is cheaper than the cleanup.

The one-line version: a single operator on a mid-tier plan is paying for a very good assistant. Three top-tier plans run in parallel pay for a staffed engineering and operations team that works around the clock, and that team's output is the product.

PART 3 · THE LANE MODEL

One lane per surface

Every terminal session maps to exactly one area of responsibility, which I call a lane. A lane owns its surface the way a small team would: its backlog, its branches, its tests and its releases.



Why lanes instead of one large session

A single session that tries to hold a whole product in its head loses the thread once its context fills, and it starts making decisions that contradict the ones it made an hour earlier. A lane holds one module, so its context stays relevant and its decisions stay consistent. It also mirrors how you would staff a human team, which makes the rest of the system natural: each lane has a backlog, a lane that finishes work writes a handoff note, and a fresh session picks the note up cold the next morning.

The lane boundaries in my setup follow the product's own modules, which are also the projects in the ticket system. That alignment matters more than it looks. When a defect appears on the billing page, there is exactly one lane it belongs to and exactly one project the ticket goes in, so nothing has to be negotiated before work starts.

The hygiene lane deserves a separate mention. It owns the system itself: the standing instructions, the memory, the hooks, and the command line tools. Every other lane reports friction to it rather than patching the shared rules mid-task. Without it, thirty sessions would each edit the shared rules to suit their own work, and the rules would drift apart within a week.

What keeps thirty sessions from colliding

Parallel sessions working in the same repository will overwrite each other unless you design against it. Each of these rules came from a collision that actually happened.

Rule	How it works	The collision behind it
Never edit the shared checkout	Every unit of work starts in its own git worktree on a fresh branch and finishes with a pull request, a squash merge and teardown in one motion.	Two sessions in one checkout shared a staging area, and one committed the other's half-finished work.
Claim before you touch	A small claims tool records which session holds which branch or file. Every session checks it before editing, merging or deploying.	The same security fix was built twice in one hour by lanes that were not talking to each other.
Serialize merges	One merge, one staging verdict, then the next. Lanes wait for an idle window before merging.	A queued test run was silently replaced by a later merge, so a change reached production untested.
Land small and often	Small pull requests, merged as soon as they are verified.	Five small pull requests merged with zero conflicts; two 133-file pull requests produced eight.
Announce production	A production release claims the main branch with a stated takeback time, so every other lane sees it in flight.	A claim held by a session that had died blocked two lanes for hours, which is why every claim now carries a takeback time.

The failure to design against is the silent one

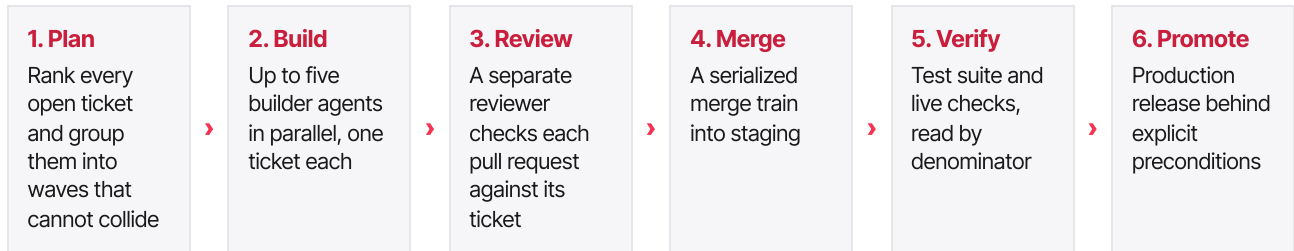
Loud failures take care of themselves. A merge conflict stops you, a failing test turns red, a broken page gets noticed. The failures that cost the most in a parallel system report success. The clearest example from my own setup: the continuous integration service keeps one pending run per queue, so when three lanes merged in quick succession, the middle lane's test run was evicted before it started. Its change showed as merged, its deploy step showed as complete, and its test suite never executed. Nothing turned red.

The rule that came out of it, merge only into an idle window and watch your own run to a verdict, costs a few minutes per merge. The alternative costs a production incident that nobody can trace, because every signal said it was fine.

PART 4 · THE BUILD LOOP

How a backlog becomes shipped code

The largest single consumer of capacity is what I call a waves run: a lane takes an entire project from the ticket system and works it to production with minimal input from me.



Planning is where the parallelism is won or lost. Two tickets that edit the same file cannot run in the same wave, and a ticket that depends on another's schema change waits for the wave after it. A good plan puts five genuinely independent tickets side by side; a careless one puts five tickets on a collision course and spends the saved time resolving conflicts.

Each builder takes one ticket from first read to an open pull request. It reads the relevant code, calls for a second opinion from a different model before its first edit, writes the change, runs the tests and the linters, and opens the pull request with its evidence attached. It never merges and never deploys to production. A reviewer on a different model then reads the diff against the ticket and against the evidence the builder claims to have produced. An overseer reviews the finished wave as a whole and steers the plan for the next one.

A supervisor process restarts the loop if a usage limit stalls it, so a wave that pauses at two in the morning resumes when capacity returns. Across the life of the ticket system, 1,729 tickets have reached Done, counted in full on 18 September 2026.

The agent roster

Agent	Job	Posture
Ticket builder	Takes one filed ticket from first read to an open pull request. Never merges, never deploys to production.	Most capable model, with a mandatory second opinion before its first edit
Ticket reviewer	Reads the pull request against the ticket and against the evidence the builder produced.	Read-only, on a different model from the builder
Waves overseer	Reviews a whole finished wave and adjusts the next one.	Read-only, and never reviews a wave it helped build

Agent	Job	Posture
Second model family	OpenAI's Codex executes tickets that are fully specified, and gives an adversarial second opinion on hard calls.	A separate family catches what the first one is blind to
Research sweep	Answers one scoped question with a source for every claim. Several run in parallel.	Read-only; reports and never decides
Inbox, CRM, meetings	Triages the inbox, audits CRM hygiene, and turns transcripts into decisions and owners.	Reads and proposes; a human applies

Two rules that make delegation safe

The executor that wrote the code never reviews it. A model reviewing its own work reproduces its own blind spots with great confidence. Review always goes to a different agent, and for the highest-risk changes to a different model family.

Anything touching authorization, billing or tenancy gets an independent review before merge, whoever wrote it. That includes code I wrote myself in a session. These are the three surfaces where a defect is invisible in normal use and catastrophic when it surfaces, because it shows one customer another customer's data or charges the wrong account.

Using a second model family

Codex from OpenAI is a first-class executor in the system, with a narrow brief. It takes a ticket only when the ticket is fully specified: the target state is stated so precisely that two competent readers could not produce different implementations, the files in scope are named, acceptance can be checked from the diff alone, and no question in the ticket or its comments is still open. On the proving run, it took three real tickets and delivered three pull requests that were independently reviewed and merged. On one of them it fixed the branch next to the one the ticket named, and listed the other callers it had checked.

The same proving run produced the most useful lesson in this part. One of the three runs stopped to ask for design approval, wrote zero files, and exited with a success code. Any dispatcher that trusted the exit code would have recorded a finished ticket that did not exist. The dispatcher now checks the file count and treats a run that wrote nothing as a run that needs a human.

The binding constraint is specification

When I tried to delegate more work to the second model, I screened the backlog for tickets that looked ready because they carried an acceptance section. Twelve did. Reading them closely, about three actually were. The rest had acceptance criteria that were open design questions wearing a heading, such as "something runs on a schedule", or comments still arguing about the mechanism.

The limit on delegation is how well the work is written, and agent capacity comes second. If you want to delegate more, the investment that pays is writing better tickets, and loosening the bar costs more than it saves.

A ticket is ready to delegate when: the target state is stated precisely enough that two readers would build the same thing; the files in scope are named or obvious from the path; a reviewer could check acceptance from the diff without reading the surrounding code; nothing in the body or the comments is still being decided; and it does not touch permissions, billing, schema or data migration. If writing the task requires reading the implementation first, the ticket is not ready, and that is useful information in itself.

Verification means running something

Every lane follows one rule about evidence: a claim is checked against something that ran, such as a test, a query, a rendered page or an API response, and never against the session's memory of what it did. A green check counts only once you know its denominator. A test suite that passes with 40 cases when it used to run 400 has told you nothing, and it looks identical to one that passed all 400.

PART 5 · THE TOOLCHAIN

The plumbing under the sessions

The model is half the system. The other half is the plumbing that lets thirty sessions act on real tools, remember what they learned, and stay inside guardrails.

Layer	What it is	Why it matters
Interface	Claude Code in the terminal for build and operations work; chat for thinking and drafting	The terminal is where parallel, long-running, file-based work happens.
Tool access	About 30 small command line tools, each wrapping one vendor: the ticket system, chat, email, documents, CRM, database, job runner, payments, accounting, workflow automation and enrichment vendors	A command line tool costs nothing until it is called. A connected tool server loads its full set of definitions into every session, whether that session uses it or not.
Skills	About 100 packaged workflows: ticket pickup, waves, proposals, pitch decks, enrichment runs, and branded PDF deliverables like this one	A repeatable process costs one skill load instead of a long explanation every time.
Memory	File-based memory with a short root index and on-demand sub-indexes, plus semantic search over roughly 1,100 memory files	Sessions learn from each other's mistakes without loading everything at boot.
Hooks	About 30 scripts that fire on session events: log every tool call, upload transcripts, announce open decisions, warn when standing context grows	Enforcement and observability that do not depend on the model remembering.
Secrets	A password manager as the only store, with credentials delivered to a tool at the moment of the call	A session can use a key without being able to read it.
Execution	A durable job runner for scheduled and long-running work, and an agent control plane for operational agents	Recurring work runs off the interactive plans entirely.
Knowledge	A knowledge graph that every meeting, client and vendor flows into	Agents answer from the company's own record instead of researching it again.

Command line tools over tool servers

This is the single most effective change I made to capacity, and it is counterintuitive. Tool servers are the default way to connect an assistant to other software, and they are convenient. Their cost is that every connected

server loads its tool definitions into every session at boot. Ten servers with twenty tools each put hundreds of definitions into every conversation, including the thirty conversations that will never touch most of them.

A command line tool loads nothing until it is called. The session knows the tool exists from one line in a skill, and it reads the help text only when it needs it. My rule is a decision tree: if a vendor ships an official command line tool, wrap it in a skill; if it has a public API, build a thin tool against it; if it only offers a tool server, build a small client for that same endpoint; and connect a tool server only when none of those works, with the reason written down.

Memory that scales past one conversation

Thirty sessions produce a lot of learning, and almost all of it would be lost without a place to put it. Each lesson becomes a small file holding one fact, with a type (a user preference, a correction, a project state, or a pointer to a resource) and a one-line description. A short root index loads at every session start and holds only what a session must know before it would think to search: the hard rules and the current state of live work. Everything else sits in sub-indexes and is found by semantic search when a session needs it.

The root index has the same budget discipline as the standing instructions. A row belongs in it only if a session needs it unprompted, not knowing it leads to the wrong action, search would not surface it, and it is still current. Anything that fails a test moves down a tier.

Hooks enforce what instructions only request

An instruction in a prompt is a request, and a model under pressure will occasionally skip it. A hook is a script the harness runs on a session event, whether or not the model remembers. Mine log every tool call to a database so any agent's work can be audited afterward, announce the open decisions waiting on me at every session start, measure the standing context against its budget, and warn when a session is about to work where another holds a claim. Most of them warn rather than block, because a false positive that wedges a session mid-task costs more than the collision it prevents. The ones that earn a hook are the failures that happen silently.

Secrets stay out of the conversation

Credentials never pass through a session's context. They live in one store, reach a tool at the moment it runs, and never appear in a transcript. Anything that could print a credential goes through a redactor that works by field name, since a redactor that matches value patterns only catches the formats you predicted, and a new vendor's key format will slip past it.

Durable work belongs off the interactive session

A scheduled job, a long batch, or a retry loop does not belong in a terminal session that ends when the laptop closes. That work runs on a durable job runner, and the operational agents run on a control plane. The interactive sessions build and supervise that work; they do not host it.

PART 6 · GUARDRAILS

What the guardrails cost, and what they prevent

Anyone running agents at scale has been burned. The guardrails in this system are the scar tissue, and each one is here because of a specific failure.

Destructive operations are denied in configuration

The operations that cannot be undone are blocked in the harness configuration, and no instruction in a prompt can override that. The classes are: destroying an environment, overwriting production data from a backup or another environment, running arbitrary code against production, and writing hand-built database changes straight into production. A session that believes one of these is necessary stops and hands it to me.

What remains allowed is deliberate. Deploying code to production through the reviewed path is routine. Writing production data is allowed only through a committed migration script that meets every one of these conditions:

- It was reviewed and merged before it runs, so nobody writes and runs it in one motion.
- It dry-runs by default and needs an explicit flag to apply.
- It captures a rollback snapshot before any write and prints the undo command.
- It is idempotent, so running it twice changes nothing the second time.
- It reports its full denominator: changed, skipped, failed, and accounted for out of the total.
- It refuses to continue on a partial read.
- It verifies its writes by reading a sample back, and never trusts the write call's own count.
- It names the person who authorized the run, and records that name in the output and the snapshot.

No agent sends email

Every message an agent writes lands as a draft, and I press send. That rule has no exceptions, including resends of messages that appear to have failed. It exists because a session once re-issued a formal notice it could not see had already gone out, and the customer received it twice. Email cannot be recalled, so the only safe design is one where an agent cannot send at all.

The decisions register

With thirty sessions running, the most expensive thing an agent can do is block on me without my knowing it. Anything waiting on my decision becomes its own file in a register, and every entry carries four things: the question, a recommendation with its reason, the context needed to decide written inline, and what happens by default if I say nothing. The open titles print at the start of every session. A default keeps work moving while I am busy, and it never counts as consent for anything irreversible or outward-facing; those wait for an actual answer.

Two failure shapes worth your time

A check that cannot fail. A guard that sits behind a branch the normal state never enters is a comment with an exit code. It passes forever, it looks like protection, and it catches nothing. Before believing any green result, I now ask what a red one would look like and whether it could actually happen.

A count read without its denominator. "Zero errors" on an empty input, "all tests passed" on a suite that silently skipped half its cases, and a list of two failures from a system that reported a hundred and twenty all look like good news. Every count in the system is reported beside what it was counted out of, and a report that cannot say its denominator does not count as evidence.

Guardrail	Failure it prevents
Destructive operations denied in configuration	An environment or its data destroyed by a session that was sure it was right
Migration scripts with snapshot, dry run and read-back	A bulk production write that cannot be undone or verified
Drafts only for email	A duplicate or premature message reaching a customer
Independent review for authorization, billing and tenancy	A defect that shows one customer another's data, or charges the wrong account
Decisions register with defaults	Thirty sessions silently blocked on a question nobody asked clearly
Every count reported with its denominator	A green result that measured nothing

PART 7 · RUNNING THE COMPANY

Agentic operations beyond the code

Engineering is where this system started, and operations is where most of a go-to-market team's time actually goes. The same pattern runs both.

The architecture

A control plane runs more than 45 agents organized into departments: revenue, delivery, growth, intelligence and operations. The design splits four responsibilities that are usually tangled together.

Responsibility	Where it lives	Why it is separate
Deciding	The agent control plane, which holds each agent's role, goals and budget	Judgment is the part that needs a model; everything around it should be plain software.
Executing	A durable job runner with retries, schedules and fan-out	Work that must finish should survive a crash, a timeout or a closed laptop.
Recording	The ticket system for work, and the CRM for customers	One record per kind of truth, so no agent keeps its own shadow copy.
Auditing	Every tool call logged to a database	When an agent's output looks wrong, you can see exactly what it did and in what order.

Defined agents for defined jobs

Six recurring go-to-market motions each have an agent with a defined role: meeting prep, new account pursuit, deal review, client implementation, campaign launch and competitive response. Each takes a specific input, such as a meeting on the calendar or an account name, and produces a specific output in a known place, such as a prep brief, a pursuit plan or a battlecard. None of them sends anything to a customer; every outward-facing result is a draft for a human.

The discipline that makes these reliable is the same one that makes the build loop reliable. Each agent has one job, a written definition, and an output someone checks. An agent with a vague brief produces plausible work that nobody can evaluate, which is worse than no agent at all.

Knowledge that compounds

Every meeting transcript, client document, vendor change and research finding flows into a knowledge graph. Before an entity is created, the ingestion step checks whether it already exists, so the graph grows without duplicating itself. The agents answer from that graph first. A pursuit agent working a new account starts from

everything the company already knows about it, including past conversations, related accounts and competitive history, and researches only the gaps.

This is where the system's value compounds over time. A new hire forgets most of what they are told in their first month. A knowledge graph keeps all of it, and every agent reads it on every run.

PART 8 · THE PERSONAL OS

The same pattern runs the house

Once the pattern worked for the company, applying it at home was a small step. A household has a backlog, recurring operations, a calendar, suppliers and a budget, and most of it is repeatable.

A private repository is the single source of truth for the household: the people and their constraints, the weekly rhythm, vendors and service providers, home maintenance, seasonal yard work and a chore system. Agents act on it through the same guarded tooling the company uses, with the same rules about drafts, credentials and destructive actions.

Module	What the agent does
Calendar	Reads and writes the shared family calendar, checks for duplicates before adding anything, and respects schedules that arrive from outside feeds.
Meals and groceries	Once a week, plans the meals against everyone's dietary rules, merges the ingredients with the staples list, and produces one consolidated order for approval.
Home and chores	Tracks seasonal house and yard cadences, service days and recurring pickups, and publishes the chore charts where the family actually looks.
Ingest	A scheduled job pulls documents from the family's shared drive into a database the agents can query.

The household is also where the approach is easiest to explain. Nobody at my kitchen table cares about worktrees or merge trains. They care that the grocery order reflects what everyone can eat, and that the calendar is right. The system underneath is the same one that ships software.

The principle is identical at work and at home: write the system down once, in files an agent can read, and keep the judgment.

PART 9 · ADOPTION PATH

Four weeks to running parallel

You do not need thirty sessions to get most of the value. Two or three well-organized lanes on one plan will change how much work you get through, and the path there takes about a month.

Week	Move	Outcome
1	Write your standing instructions: who you are, how you work, what good output looks like, and what is forbidden. Add a deny list in configuration for every command that cannot be undone.	Every session starts aligned, and the worst mistakes are impossible.
2	Give each project or client its own repository and its own terminal lane. Start every change in a git worktree and finish it with a pull request.	Two or three sessions running without collisions.
3	Turn your three most repeated workflows into skills. Replace any daily tool server with a small command line tool if one exists.	Less repeated prompting, and lighter sessions.
4	Put work in a ticket system, add a builder and a reviewer agent, and let one lane work a backlog while you work in another.	Your first autonomous loop, and the point where a larger plan pays for itself.

Five habits that stretch any plan

1. **Clear context between tasks.** Write a short handoff file, then start fresh. A long tail of stale conversation is paid for on every turn.
2. **Keep standing instructions lean.** Move history and evidence into reference files that load on demand, and measure what loads at boot.
3. **Match the model to the job.** The most capable model for design, debugging, and anything touching money or permissions; lighter models for mechanical edits where a mistake would fail loudly.
4. **Pass paths, not file contents,** when you delegate to a subagent, and have it write its summary to a file.
5. **Specify work precisely.** The more exactly a task is written, the less the agent explores, and exploration is where capacity disappears.

What the first week feels like

Slower, and you should expect that. Writing standing instructions, setting up worktrees and learning to specify a ticket properly all take time that a single chat session does not. Most people quit in that first week because the old way still feels faster.

The payback arrives the first time you come back from a meeting, or wake up, and find a lane has worked through a backlog without you: tickets built, reviewed, merged and verified, with a handoff note waiting and a short list of decisions only you can make. From that point, the question changes from how much you can do in a day to how much you can supervise.

Where people usually get stuck

Sticking point	What usually fixes it
Sessions overwrite each other	Worktrees for every change, and a claims check before any edit or merge.
Agents keep asking the same questions	The answer belongs in the standing instructions or in memory, written once.
Delegated work comes back wrong	The ticket was underspecified. Rewrite it until two readers would build the same thing.
Limits arrive early in the week	Standing context has grown, or subagents are being handed file contents instead of paths.
You cannot tell what is waiting on you	A decisions register, with a recommendation and a default on every entry.

PART 10 · THE CLOSE

Everything is in this document

Nothing was held back. Every part of the system that runs GTMify is described in these pages, in enough detail to build it yourself.

Most readers will change nothing, and that is the honest outcome. Building this took months, and every rule in it came from a failure I had to live through first. If you take two habits from Part 9, you have already gotten your time back.

If you want the result without spending the months, there are two ways to get it.

Path	What it means
Have it installed	We set up the operating system in your business: standing instructions, lanes, guardrails, the builder and reviewer loop, and the first skills, scoped to your work.
Use what it produces	GTMify is the go-to-market product this system builds and runs. If what you need is the revenue engine rather than the machine that builds it, start there.

Either way, the first step is the same: a 20-minute parallel audit. We look at how you work today, find where your attention is the bottleneck, and tell you which of the two paths fits, or whether you need neither.

[Book the 20-minute parallel audit](#)

Scott Wueschinski, Co-Founder, GTMify.io · September 2026