

When Should I Arrive at the Airport? A Newsvendor Perspective on the "Airport Dad" Debate



GAD ALLON
OCT 06, 2025



9



1



Share



This Week's Focus: The Myth of the Three-Hour Rule

Airlines urge travelers to arrive two to three hours before departure—a rule few will question and a habit that costs U.S. travelers an estimated \$83 billion a year in lost time. Yet data show that the average TSA waits are under 20 minutes at most airports. This week, we treat the “airport theory” like a newsvendor problem—balancing the peace of mind of arriving early against the opportunity cost of wasted time—to ask: Has “airport Dad” been right all along?

Airlines often advise travelers to be at the airport **two to three hours** before departure—a rule of thumb so ingrained that it’s rarely questioned.

This habit, however, proves expensive. According to Gary Leff—an expert in the field

of miles, points, and frequent business travel—this requirement is a “**system failure no one’s trying to fix**.” He argues that longer dwell times have turned airports into shopping malls, and notes that airlines and airports both share in the revenue generated from captive passengers. In 2024, U.S. airlines carried about **862.8 million domestic passengers**. “Airline passengers skew higher income, so let’s conservatively assume a \$100,000 average income or \$48.08 hourly wage. Taking an extra two hours per passenger on average, that’s **1.725 billion hours, or \$83 billion cost to the economy** just for extra time wasted for domestic passengers.” And this figure, he notes, only considers the outbound wait, not the additional hours travelers spend waiting for luggage, buses to rideshare lots, or connecting flights.

Indeed, different studies show data indicating that the “airport theory”—the belief that one must arrive several hours early—may be unnecessary. A 2025 analysis of security wait times at **31 major U.S. airports** found that the longest average wait was **21 minutes** (Honolulu); Denver and Chicago O’Hare averaged **one minute**!

Similar studies in Europe are even more encouraging: at **Paris-Charles de Gaulle**, **94.5 %** of passengers cleared border controls in **under ten minutes**, while most British airports reported average security waits between **10 minutes and 20 minutes**. These findings suggest that many travelers may be unnecessarily sacrificing valuable time.

Nate Silver, known for political forecasting, deserves credit for attempting to operationalize the tradeoff many people face: arrive too early and spend too long at the lounge; arrive too late and possibly miss your flight. He created a spreadsheet to help people decide when to leave for the airport, building in adequate time buffers.

This issue has all the hallmarks of a newsvendor problem: balancing overage and underage. And the goal of this newsletter has always been to add a level of rigor, both to the tradeoff as well as the data, and to help my readers make better decisions.

Let’s dig deeper.

Subscribe

Newsvendor and the Flaw of Averages

When thinking about airport arrivals, it’s tempting to ask, “How long does it usually take?”

If the average commute is 20 minutes, the average TSA wait is 15, and the average walk is 10, then we conclude that 45 minutes should do the trick. This is the **flaw of averages**: plans based on average inputs are, on average,...wrong.

The airport is a textbook case of the **newsvendor problem**, a centuries-old model in operations research. A newsvendor has to decide how many newspapers to stock each day without knowing the exact demand. If she stocks to the *average demand*, she will routinely run out (stockouts) or be stuck with excess inventory. The rational solution is not to aim for the average, but to choose a **quantile** of the demand distribution, determined by the relative cost of overstocking versus understocking.

Airports work the same way.

Airports work the same way:

- The “inventory” is a time **buffer**.
- A “stockout” is **missing the flight**.
- The overage cost is the **opportunity cost of waiting**.
- The underage cost is the **catastrophic loss of a missed flight**.

The math is simple (yet powerful):

$$F_S(L^*) = \frac{C_o}{C_o + C_u}$$

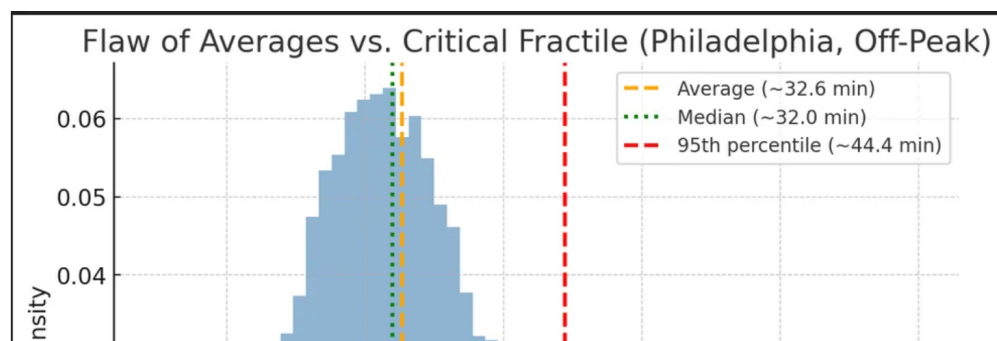
Here, F_S is the cumulative distribution of your curb-to-gate sojourn time, C_o is the cost per minute of waiting too long, and C_u is the one-time penalty of being late. The optimal arrival buffer L^* is the quantile of the distribution that balances those costs.

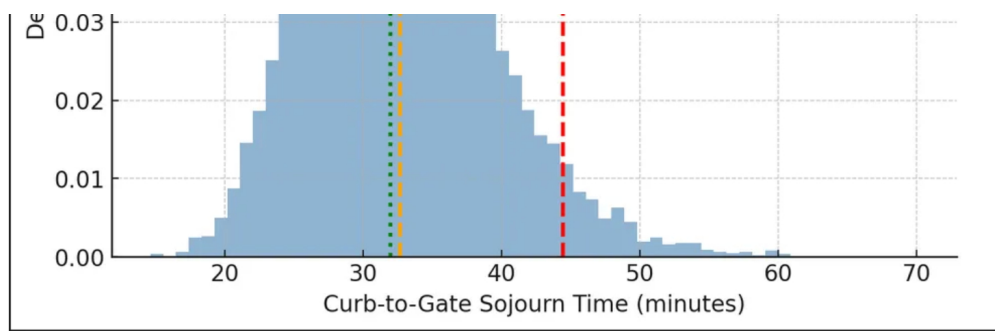
This is why **averages are misleading**. If the average TSA line is 15 minutes, but the 95th percentile is 45 minutes, your optimal buffer is determined by that 95th percentile—not by the mean. If you’re a business traveler with thousands of dollars riding on a meeting, you optimize to the 98th or 99th percentile, which may mean arriving an hour earlier than the average would suggest. Conversely, if you’re a leisure traveler who can tolerate rebooking, you might rationally operate closer to the 90th percentile and reclaim 30 minutes of your life.

The flaw of averages is not unique to airports. It’s why supply chains carry safety stock, banks stress-test against tail events, and insurance companies model 100-year floods instead of average rainfall. But airports make the lesson concrete. Asking “What’s the average wait at TSA?” is not the right question. The right question is “**What’s my cost of being early versus being late, and what quantile of the distribution does that imply?**”

The “two hours before departure” rule survives not because it’s efficient on average, but because it’s a **blunt 99th percentile hedge** that works tolerably well across travelers and times of day. The challenge—and the opportunity—is to stop planning for the average and start planning against the tails.

The following graph illustrates these points for a trip from Walnut Street in Philadelphia to Philadelphia International Airport:





The crux of the newsvendor model is the costs of overage and underage.

Underage Cost—Penalty for missing a flight: What happens if you’re late? Most airlines treat a no-show as a forfeited ticket; to rebook, you must purchase a new ticket or pay a change fee. According to a travel-agency analysis, U.S. legacy carriers typically charge about **US\$200** to change a basic economy ticket, in addition to any fare difference. Many airlines also impose a **no-show fee** in the range of **\$100–\$200**. Because missing a flight can also mean additional hotel nights and lost time at your destination, some travelers might value the penalty even higher—perhaps \$500 or \$1,000—especially on international trips or when catching a cruise. Therefore, the **cost of being late** in our model varies widely according to individual circumstances. I fly a lot (only half my flights are with United, and I still achieved Global Services status), and my travel plans are fairly optimized. Missing a flight can easily cost me between \$1,000 and \$5,000 (even \$10K in certain cases).

Cost of waiting: The cost of arriving early is more straightforward. We value travelers’ time at the mean income for U.S. domestic flyers—we use the numbers from Gary Leff’s article. At **\$100,000** average annual income ($\approx \$48.08$ per hour), waiting an extra minute costs **\$0.80** ($48.08/60$). Some travelers enjoy spending time in the lounge or browsing shops, reducing the perceived cost of waiting, while business travelers may value their time more highly. Airlines and credit cards are investing in better lounges, and in my case, many of my weekly articles have been written at lounges. I view them as an extension of my office, and my cost of being there is very low.

The point is that these costs are highly personal and the newsvendor framework accommodates these differences by allowing passengers to set their personal **cost of waiting** and **cost of being late**.

Share

The Curb-To-Gate Problem

The real decision isn’t “Do I like sitting at Hudson News for 90 minutes?” but “How much safety stock (buffer time) do I need to protect myself from stockouts (a.k.a. missing my flight)?”

We can formalize the passenger journey as a **sojourn time** random variable:

$$S = T_{\text{commute}} + T_{\text{TSA}} + T_{\text{walk}} + T_{\text{misc}}$$

- T_{commute} : the stochastic drive, drop-off, or parking shuttle. Congested at 6am Mondays, light at 11pm Tuesdays.
- T_{TSA} : a queueing problem par excellence—gamma or ex-Gaussian service times, with fat tails when an X-ray machine fails or alarm rates spike.
- T_{walk} : distance-driven, but with discrete jumps (some terminals require waiting for an APM or train).
- T_{misc} : ID check hiccups, random screening, boarding pass scans.

The deadline is not “scheduled departure” but **gate-closure** (typically 15 minutes before domestic flights, 30–60 minutes before international flights). If you might need counter service, United even imposes a **45-minute check-in cutoff**.

How do you actually know your Sojourn Time (S)?

Security distributions. The MyTSA app, live airport feeds, and studies like Transfeero/Qsensor give medians and means. The key is variance: fat tails matter more than averages. TSA operations papers show ex-Gaussian and gamma fits with heavy tails.

Walking time. Different studies map walking distance distributions; add a uniform distribution for train headways—in Atlanta or Dallas, missing a train can add 10 minutes.

Commute. Google Maps historical traffic data by day/hour gives a lognormal-ish distribution, often bimodal (rush vs. off-peak).

Two refinements matter for rigor:

Dependence. Commute time and TSA wait aren’t independent—morning rush hits both. Use copulas or bootstrapping to model joint distributions.

Robustness. At your **home airport**, you can estimate F_S from repeated experience. At **away airports**, you face parameter uncertainty. The robust newsvendor literature says: inflate your quantile by a conservatism margin. Translation: add 15–20 minutes if you’ve never flown through Dallas in August.

Calibrating costs gives intuition. Suppose your productive time is worth \$25/hour in the lounge and missing the flight costs \$1,000. Then your critical fractile is ~99%. If the 99th percentile of your airport’s sojourn time is 65 minutes, you arrive at the curbside 80 minutes before gate closure. At a familiar airport, that’s rational. At an unfamiliar one, tack on a robustness premium.

If you enjoy \$17 hummus bowls at PHL, set C_0 and arrive yesterday. If you’d rather shave 20 minutes off your buffer, buy PreCheck, CLEAR, or a ticket at Changi—variance reducers that let you lower safety stock without raising stockout risk.

Rush Hour in Philadelphia: A Distributional View

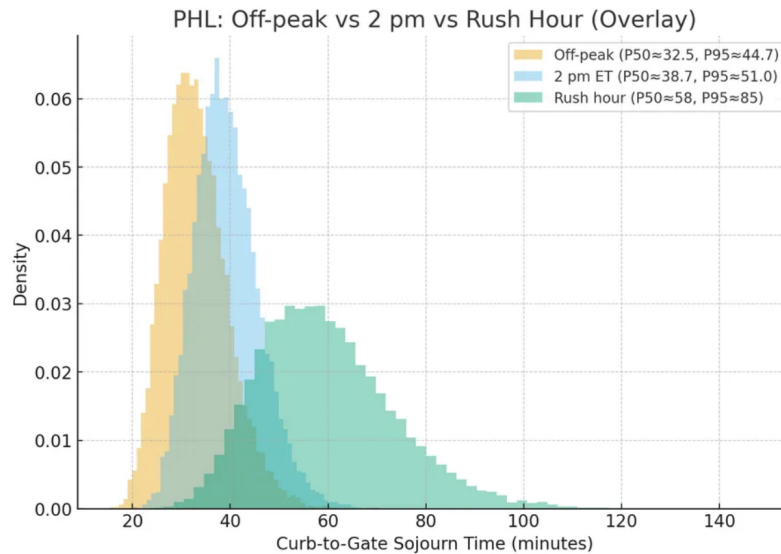
To illustrate this logic and that we can all do this math at home, I decided to measure this for **Philadelphia International (PHI)** by combining the following two data feeds:

times for Philadelphia International (PHL) by combining the following two data feeds.

Commute times: pulled from the Google Maps Distance Matrix API, which provides live, traffic-aware travel times from a fixed city origin to the airport.

TSA waits: pulled from the MyTSA feed, specifically the most recent PreCheck line wait times reported at PHL.

These were augmented by a constant walking time (10 minutes) and a small allowance for miscellaneous checks (2 minutes). The result is an empirical distribution of total curb-to-gate sojourn times, sampled repeatedly over a short collection window.



During the off-peak, you have a ~13-minute commute, a ~7-minute TSA PreCheck line, plus walking and miscellaneous. The shape of the distribution was slightly right-skewed (a few longer samples from commute variance), but with no heavy tails. At this time of day, the newsvendor model (with a moderate opportunity cost vs. a large miss cost) suggested a **leave-home time of ~53 minutes before scheduled departure**. That's astonishingly lean compared to the “two-hour” heuristic.

Rush hour does not fundamentally change the distribution's *shape*—it still looks like a lognormal-plus-gamma mixture, with a bulk of mass in the 40–70 minute range and a right tail. What it does is **shift the distribution to the right**: The commute, which had a mode near 13 minutes off-peak, stretches into **25–35 minutes with congestion**. The TSA line, which had a median of around 7 minutes with PreCheck, lengthens into **20–25 minutes**, with occasional excursions to **40+ when throughput falls**. Walking and miscellaneous are roughly fixed, so the variance is driven by commute and TSA. Under these rush-hour inputs, the **sojourn distribution's quantiles inflate**: **Median (P50)**: ~55–60 minutes, **P95**: ~80–85 minutes, and **P99**: ~95–100 minutes.

The critical-fractile rule, applied to this distribution, gives different answers depending on traveler type:

Leisure traveler (90% service level): ~75 minutes curb-to-gate. With a 15-minute gate-closure buffer, that's **90 minutes before departure**.

Business traveler (95–98%): ~80–85 minutes curb-to-gate. With a buffer, that's **100 minutes** before departure.

Risk-averse/families (99%): ~95–100 minutes curb-to-gate. With a buffer, that's **2 hours** before departure.

So the “two-hour” rule of thumb (albeit calculated from the time one leaves for the airport) is not exaggerated—it corresponds to the **99th percentile** of the empirical rush-hour distribution. The real waste only happens when we are forced to apply this rule universally, even during off-peak hours when the distribution is centered 30 minutes to the left.

[Subscribe](#)

Making Sense of the Results

Generalizing the results from PHL:

Rational behavior is driven by distributions, not averages. A median of 39 minutes (at 2 pm) says little about reliability if the 95th percentile is 51 minutes and the 99th percentile is 56. Rational travelers optimize to these higher quantiles, not the mean, because the cost of missing a flight dwarfs the cost of waiting.

A median of 55 minutes is irrelevant if your personal cost of missing the flight is thousands of dollars.

Variance is policy-driven. Airlines and TSA can shrink buffers by shrinking variance—better staffing, transparent wait-time feeds, and faster security technology.

Traveler heterogeneity matters. At PHL, a PreCheck business traveler departing at midday and targeting the 95th percentile needs ~66 minutes before departure (51mins curb-to-gate plus 15mins for gate closure). A family without PreCheck at morning rush hour, facing longer TSA and commute tails, may rationally need closer to 2 hours.

The math explains the “wasted” time as safety stock. The real focus isn't to shame buffers, but to **make the distribution tighter** so that the 95th percentile converges toward the median—and the rational recommendation shrinks naturally.

But there are still scenarios where leaving earlier is best:

Uncertain or exceptional circumstances: holidays, strikes, government shutdowns, severe weather, system outages, or staffing shortages can lengthen queues. Even at well-run airports such as Paris, a breakdown in the PARAFE system temporarily pushed the average border-control waits to **45–50 minutes**. When such disruptions are likely, conservative arrival times make sense.

Long-haul international flights: Missing an intercontinental flight or a cruise may impose high penalties (hotel rebooking, lost prepaid tours). In our model, this translates to a high C_u . Even so, the optimal arrival time rarely exceeds **two hours** in Europe and is barely above **1 h 40 m** in the U.S.

Families or first-time travelers: People unfamiliar with airports, those travelling with young children, or those requiring assistance may prefer to arrive early to avoid stress. The cost of waiting is lower for holidaymakers who treat the airport as part of the vacation.

Peak congestion: Some U.S. airports experience specific peaks. For example, BWI’s average wait is **10 minutes**, but at **8 p.m. on Sundays** it jumps to about **13 minutes**. Monitoring peak times via apps or TSA checkpoints can help adjust arrival times.

In every case, the newsvendor model is only as good as its inputs. Real-life complexities—parking, bag-drop lines, unpredictable road traffic, and airports built like mazes—add variability. Some travelers value peace of mind more than an hour of additional sleep.

In such cases, the “airport dad” may win the argument not because the data support a three-hour early arrival, but because stress and social factors matter. Still, the model suggests that **a buffer of two or more hours is rarely optimal**.

Bottom Line

The flaw of averages means two hours might actually be the right decision. And, while wasteful, unless TSA, airlines, and city traffic engineers shrink the variance, the safest strategy will remain exactly what it has been for a century of operations research: carry more safety stock, or in this case, arrive embarrassingly early.

Leave a comment

Thanks for reading Gad’s Newsletter! Subscribe for free to receive new posts and support my work.

Type your email...Subscribe



9

1

Share

← Previous

Discussion about this post

CommentsRestacks

Write a comment...



Alan Malter 8h

...

I love the topic, which is probably near and dear to the hearts of most of your readers. Your analysis raises lots of related questions. One I had was only mentioned at the very end -- the cost of stress vs value of peace of mind, which may boil down to interpersonal differences. Some travelers are stimulated by cutting it close but still making their flights, while others are strict stress-avoiders, all else being equal. Other related cost factors include whether the flight is the last one of the day (on that route) or whether you are traveling alone or with a group/family; in both cases, rebooking a missed flight becomes significantly more complicated. And if you are interested in delving deeper into this topic, what is your wisdom on minimum/optimal layover times for connecting flights? Online reservation systems often recommend flights with suspiciously short connection times! It's one thing to make it in time to the first leg of your trip, but if you miss the connecting flight then you're still in a pickle (or worse, because checked luggage successfully reaching the second flight in time is a separate problem). Finally, your essay is missing important empirical context on the real extent of this issue: Most people worry about missing flights, but how many passengers actually miss flights, and for what reasons? (i.e., are there data on what percentage of passengers actually miss flights because they did not allow enough time to get to the airport/to the gate?).

♡ LIKE 💬 REPLY

📤 SHARE

Top Latest Discussions



Get Off My Sidewalks

We all want fast delivery ... food, groceries, products...

JUL 18, 2022 • GAD ALLON



The Case Against Self-Checkouts

The Associated Press News reported on an interesting poll this week showing that 70% of people prefer a hybrid store that allows you to choose between...

AUG 23, 2021 • GAD ALLON



Southwest Airlines Meltdown: The Real Root Cause

If you follow the news, you've probably read the many articles written about Southwest Airlines' meltdown during the days around Christmas:

JAN 2, 2023 • GAD ALLON



See all >

Ready for more?

Type your email...

Subscribe

© 2025 Gad Allon • [Privacy](#) • [Terms](#) • [Collection notice](#)

Start writing

Get the app

Substack is the home for great culture

