

Lotsen würfeln nicht – oder doch?

Systemische Herausforderungen und Werkzeuge für robuste Entscheidungen im probabilistisch geprägten Luftverkehr

Jörg Buxbaum, DFS Deutsche Flugsicherung GmbH

Abstract

Dieser Beitrag zeigt beispielhaft, warum operative Flugsicherung im Flughafennahbereich und bei gemischter Pistenutzung für Starts und Landungen als **streuungsgetriebenes System** verstanden werden muss. An drei praxisnahen Szenarien – **Sequenzierung im Endanflug**, **Mixed-Mode-Betrieb** und **Kommunikations-Blocking** im Gegenanflug – wird deutlich, wie sich kleine Varianzen in Reaktions-, Umsetzungs- und Kommunikationszeiten in jenen Sekundenfenstern aufschaukeln können, in denen Kapazität und Sicherheit entschieden werden. Reihenfolgeeffekte, die stochastische Unsicherheit bei Startumsetzungen (Reaktion + Pistenbelegungszeit) und temporär blockierte Funkkanäle können **emergente Auswirkungen** auf Abstandsmargen, Restflugstrecken und Landerate haben – selbst bei ansonsten stabilen Rahmenbedingungen.

Aus dieser Diagnose leitet die Arbeit einen **Perspektivwechsel** ab: von **Punktwerten** zu **Wahrscheinlichkeitsräumen**. Sicherheitsmargen werden als **Risikobudgets entlang der Kette Wahrnehmen/Entscheiden – Übertragen/Bestätigen – Verstehen/Umsetzen** geführt; **Gates** ersetzen starre Zeitpunkte; **Quantile** sind operativ bedeutsamer als Mittelwerte. Daraus entstehen konkrete Gestaltungsempfehlungen für die nächste Generation von Assistenz: Auf Vergangenheitswerten, neuen Informationsquellen und systemischen Betrachtungen basierende Quantile in bestehenden Tools **sichtbar machen**, **kurze, begründete Empfehlungen** bereitstellen, **Medienbrüche reduzieren** (z.B. Flight Management System (FMS)-ladbare Freigaben) und Empfehlungen **kalibrieren**. Damit bleibt der Fluglotse „supervisory decision maker“, agiert aber robuster, weil Entscheidungen **probabilistisch vorbereitet** und **transparent erklärbar** sind.

1. Einleitung

Im hochdynamischen und sicherheitskritischen System der Flugverkehrskontrolle stellen probabilistische Effekte eine konstitutive Herausforderung dar.

Variabilität in der Umsetzung von Freigaben, zufallsbedingte Änderungen in der Trajektorienentwicklung sowie stochastische Einflüsse durch Wetter, Verkehrssituation und menschliches Verhalten führen dazu, dass operative Realität nur in begrenztem Maße deterministisch abgebildet und vorausgesagt werden kann. [1]

Während die strategische Planung auf möglichst stabile Annahmen setzt, ist die taktische Ebene von inhärenter Unsicherheit geprägt. Diese Unsicherheit bedeutet nicht mangelnde Kompetenz oder fehlende Systemkontrolle, sondern ist im wissenschaftlichen Sinn zu verstehen: Der künftige Systemzustand ist nicht eindeutig bekannt, Entscheidungen müssen auf der Basis unvollständiger, streuungsbehafteter oder lediglich probabilistisch beschreibbarer Informationen getroffen werden.

In solchen Kontexten entscheiden Wahrscheinlichkeiten, Erfahrungswerte und systemisches Einschätzungsvermögen häufig über den tatsächlichen Handlungsspielraum.

Ziel dieses Beitrags ist es, ausgewählte probabilistische Phänomene in der Flugverkehrskontrolle darzustellen und hinsichtlich ihrer Relevanz für die Automatisierung von Entscheidungsprozessen zu analysieren.

Hierbei werden exemplarische Situationen aufgegriffen, in denen Streuungen und stochastisch geprägte Ereignisse einen unmittelbaren Einfluss auf Sicherheit, Effizienz und Kapazität haben – etwa bei

- der Optimierung von Anflugreihenfolgen,
- der Dimensionierung von Staffelpuffern oder
- der Auswirkung von Initial Calls in hochbelasteten Lufträumen.

Durch geeignete Modellierungsansätze sollen sowohl die Grenzen deterministischer Systeme aufgezeigt als auch Ansatzpunkte für robuste, probabilistisch informierte Automatisierungskonzepte identifiziert werden.

2. Begriffsrahmen und systemischer Kontext

Die differenzierte Auseinandersetzung mit probabilistischen Phänomenen in der Flugverkehrskontrolle erfordert zunächst eine präzise Begriffsklärung (*Abbildung 1*). Unter *Determinismus* [2] wird ein (philosophisches) Systemverständnis gefasst, in dem jedes

Ereignis durch vorhergehende Zustände eindeutig festgelegt ist. Deterministisches Verhalten ist eine konkrete Ausprägung dieser These auf Systemebene, die sich in der Gleichheit von Ein- und Ausgaben zeigt. Im Gegensatz dazu beschreibt der *Zufall* [3] das Auftreten von Ereignissen ohne erkennbare Kausalität oder Regelmäßigkeit.

Zwischen diesen Polen liegt die *Probabilistik* [4] (im Sinne von Wahrscheinlichkeitsmodellierung), die Prozesse beschreibt, deren Entwicklung nicht exakt vorhersagbar, jedoch durch Wahrscheinlichkeitsverteilungen quantifizierbar ist. *Höhere Gewalt* [5] beschreibt ein von außen kommendes, weder vorhersehbares noch abwendbares Ereignis, das auch durch äußerste Sorgfalt nicht verhindert werden kann. *Koinzidenz* [6] bezeichnet das gleichzeitige Eintreten mehrerer – zumeist unabhängiger – Ereignisse, das in komplexen operationellen Kontexten zu überraschenden Systemzuständen führen kann.



Abbildung 1: Begriffsrahmen - Einordnung der Begriffe im Spannungsfeld aus Vorhersagbarkeit und Beeinflussbarkeit

Diese Konzepte sind keineswegs abstrakt, sondern entfalten in der Flugsicherung eine unmittelbare praktische Relevanz. Während viele Regelwerke, Verfahrensbeschreibungen und Systementwürfe deterministische Strukturen unterstellen, zeigt sich die operative Realität häufig als probabilistisch geprägt. Bereits in alltäglichen Abläufen – etwa bei der Umsetzung von Freigaben, der Reaktion auf konfligierende Flugpfade oder der Bewertung und Abschätzung meteorologischer Entwicklungen – müssen Fluglotsen mit einer inhärenten Unsicherheit umgehen. [7] [8] [9]

Ursache ist eine Kombination aus Informationslücken, Verarbeitungsgrenzen und Strukturkomplexität: Relevante Daten sind teils nicht verfügbar oder nur unvollständig/rauschbehaftet; verfügbare Informationen können in der gegebenen Zeit weder kognitiv noch maschinell zuverlässig integriert werden. Dazu können Zusammenhänge so komplex bzw. verdeckt sein, dass daraus keine verlässlichen, operativ nutzbaren Schlussfolgerungen gezogen werden können.

In solchen Situationen genügt es nicht, auf starre Entscheidungsbäume zurückzugreifen; vielmehr sind flexible, kontextadaptive Strategien gefragt, die Wahrscheinlichkeiten und Erfahrungswissen integrieren.

Vor diesem Hintergrund lässt sich die Flugverkehrskontrolle als komplex-adaptives, soziotechnisches System charakterisieren. Es ist geprägt durch eine Vielzahl interagierender Akteure (Cockpitcrews, Fluglotsen, technische Systeme), durch sich dynamisch wandelnde Rahmenbedingungen (Verkehrslage, relevante Wetterphänomene, Infrastrukturverfügbarkeit) sowie durch emergente Effekte, bei denen aus lokal sinnvollen Einzelentscheidungen globale Systemzustände entstehen, die nicht deterministisch kausal vorhersehbar sind. [10] *Abbildung 2* zeigt exemplarisch relevante Wettereffekte und die erreichbare Prognosegüte in Abhängigkeit vom zeitlichen Vorlauf. Besonders die Nebelaufklärung am Morgen/Vormittag stellt an hochbelasteten Flughäfen eine hochrelevante Erscheinung dar, deren Vorhersage trotz intensiver Forschung weiterhin mit erheblicher Unsicherheit behaftet ist. [11] [12]

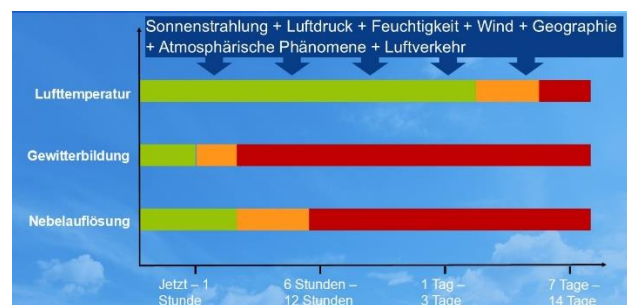


Abbildung 2: Vorhersagehorizonte und damit verbundene Genauigkeiten für verschiedene Wetterphänomene. „Grün“ für eine für ATM-Zwecke hinreichend gute Vorhersehbarkeit, „gelb“ für streuungsbehaftete Prognosezeiträume und „rot“ für Zeiträume mit hoher Unzuverlässigkeit. Eigene Darstellung auf Basis der praktischen Erfahrungen als Leiter Anflugkontrolle Frankfurt (DFS Deutsche Flugsicherung GmbH)

Für den Terminal Manoeuvring Area (TMA)- und Tower-Betrieb wirken z.B. Zonen konvektiver Strukturen als stochastische Engpassgeneratoren: Sie verschieben Trajektorien, erzwingen Vektoring, verändern Anfluggeschwindigkeiten und verbreitern die Verteilungen von Headway und Runway-Occupancy. Praktisch bedeutet das: deterministische Punktzeiten im Final sind zu diesen Zeiten zu knapp kalkuliert. Unter Wetterlast steigen die Tail-Risiken (z.B. verspätete Freigaben durch zusätzliche Funklast oder unerwartete Speed-Anpassungen), sodass Sicherheitsmargen in diesen Situationen als Risikobudgets geführt werden müssten. In dieser Logik werden Gates bewusst vorgezogen, Kommunikationsfenster als

Wahrscheinlichkeitsräume betrachtet und Mixed-Mode-Puffer kontextabhängig erhöht. Die Aussagen des Beitrags gelten damit nicht nur für „schönes Wetter“, sondern gerade für operationell kritische Tage.

Besonders deutlich wird dies in Phasen hoher Verkehrsdichte, etwa bei Go-Arounds (G/A) (Abbildung 3) oder rasch wechselnden Wettersystemen. In solchen Lagen kann ein einzelnes Ereignis – beispielsweise ein verzögerter Initial Call – eine Kaskade nachgelagerter Reaktionen auslösen und das System temporär destabilisieren. Dieser Effekt lässt sich auch mit Prinzipien der Warteschlangentheorie beschreiben: Sobald sich das System an seiner Kapazitätsgrenze bewegt, führt bereits eine kleine Verzögerung zu überproportionalen Auswirkungen.

Für das operative Management in der Flugsicherung bedeutet dies, dass gerade unter Bedingungen hoher Verkehrsdichte robuste Puffer, flexible Kapazitätsreserven und klare Kommunikationsabläufe erforderlich sind, um die Stabilität aufrechtzuerhalten.



Abbildung 3: Prinzipbild eines Go-Around-Manövers eines Anflugs im Arrival-Peak eines Flughafens mit der Folge einer temporären Überlastung in der Downwind-Sequenz

Auch auf strategischer Ebene findet dieses Verständnis zunehmend Eingang: Die ARIS-Strategie der EU [13] beschreibt das Luftverkehrsmanagement explizit als ein System, das „adaptive robotics, human-machine interaction und cybersecure automation“ benötigt, um unter komplexen, volatilen und sicherheitskritischen Bedingungen resilient agieren zu können. Ebenso hebt der SESAR ATM Master Plan 2025 [14] die Notwendigkeit hervor, durch datengetriebene, dynamisch rekonfigurierbare Luftraummodelle (z. B. *dynamic airspace configurations*) sowie durch *human-machine teaming* und *Artificial Intelligence (AI)-based trajectory prediction* besser auf Unsicherheiten reagieren zu können.

Die Herausforderung für eine belastbare Automatisierung liegt somit nicht allein in der Abbildung aller nur vorstellbaren deterministischer Abläufe und deren reibungslosen Kombination, sondern in der vernetzten Einbindung probabilistischer Systemmerkmale

und stochastischer Einflüsse in die operative Entscheidungsunterstützung.

3. Modellierungsansätze probabilistischer Effekte

3.1 Einfluss der Anflugreihenfolge auf Wartezeiten und sequenzbedingte Kosten

Die Bildung von Anflugsequenzen zählt zu den zentralen Aufgaben der taktischen Flugverkehrskontrolle im Endanflugbereich. Insbesondere in hochfrequentierten Lufträumen beeinflusst die gewählte Anflugreihenfolge unmittelbar die Länge der Anflugsequenz, die Minimierung von Staffelungspuffern und die Auslastung der maximalen Pistenkapazitäten. Dabei entspricht die tatsächliche Ankunftszeit von Anflügen in der TMA in den seltensten Fällen der im Flugplan hinterlegten Zeiten, sondern ist vielmehr das Ergebnis zahlreicher, teilweise unvorhersehbarer Ereignisse und Umstände im Betriebsablauf am Ereignistag:

Herkunfts-Flughafen	▪ Boarding-Prozesse: verspätete Passagiere (ggf. erneute Ausladung von Gepäck), zeitaufwändige Security-Checks, verspätete Crew oder verspätete Catering-/Bodenabfertigung. ▪ Technische Checks am Gate: kleinere Defekte oder verspätete Freigabe durch Technik. ▪ Außenposition statt Finger: zusätzliche Zeit fürs Busboarding. ▪ Flughafeninfrastruktur: Engpässe bei Handling, Gate-Verfügbarkeit oder Bodenverkehr. ▪ EUROCONTROL/Network Manager Maßnahmen: Air Traffic Flow Management (ATFM)-Slots oder Re-Routings
Taxiing / Take-Off	▪ Startbahnkapazität und Sequenzierung: Wartezeiten am Holding Point durch hohe Verkehrsdichte. ▪ De-Icing im Winter: verlängert Push-Back- bis Startzeit erheblich.
Enroute	▪ Umfliegen von militärischen Sperrgebieten / temporären Luftraumsperrungen. ▪ Air Traffic Control (ATC)-bedingte Level Capping oder

Lotsen würfeln nicht – oder doch?

	Umwege (z. B. wegen Sektorauslastung). ▪ Konfliktauflösung: Vektoren oder Speed-Anweisungen, um Staffellungen sicherzustellen.
Zielbereich (TMA/Anflug)	▪ Priorisierung bestimmter Flüge: z. B. medizinische Notfälle, Anschlüsse mit vielen Umsteigern (airline-interner Tausch). ▪ Runway-Änderungen kurz vor Anflug: z.B. durch Drehung der Windrichtung.

Die in *Abbildung 4* dargestellten Histogramme verdeutlichen diese betrieblichen Einflüsse am konkreten Betriebstag: Sie zeigen die Verteilung der tatsächlich realisierten Off-Block-Zeiten (blau) und On-Block-Zeiten (grün) eines Beispielumlaufs (Flughafen Frankfurt Main (FRA) – Flughafen München (MUC)). Die dünnen Vertikallinien markieren die Sollzeiten (Scheduled Time of Departure STD (blau) und Scheduled Time of Arrival STA (grün)).

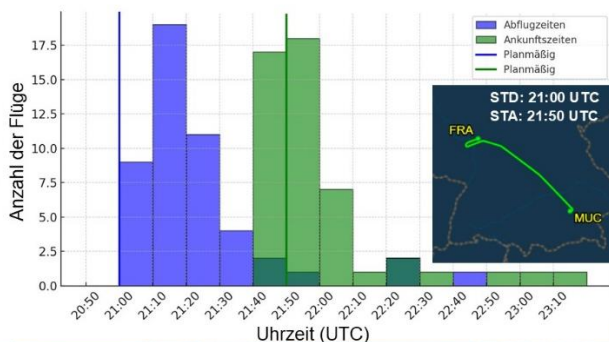


Abbildung 4: Zeitliche Lage der geplanten und tatsächlichen Off- und On-Blockzeiten von 49 Flügen München – Frankfurt (Deutsche Lufthansa [DLH]123) im September und Oktober 2024 [flightradar24]

Für die taktische Steuerung im Endanflug bedeutet das: Planzeiten sind Zufallsvariablen oder profan artikuliert: **Dass ein Flug tatsächlich auf die Minute genau wie geplant startet und landet ist Zufall.** Sequenzen und Staffelpuffer sollten daher probabilistisch ausgelegt werden (z.B. mit 80/90/95-Perzentil-Fenstern statt fester Punktwerte) und dynamisch an die beobachtete Tagesstreuung angepasst werden.

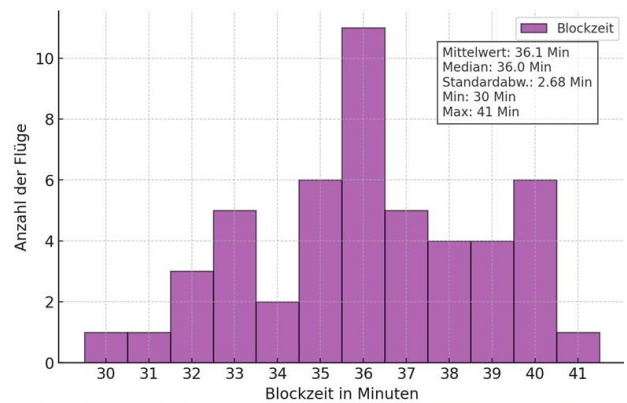


Abbildung 5: Verteilung der tatsächlichen Blockzeiten für die ausgewählten 49 Flüge der Flugnummer DLH123 im September und Oktober 2024 [flightradar24]

Eine Parallele stellt dar, dass die geplante Blockzeit (SBT) für eine Flugplansaison auch nicht aus Strecken und Flugleistungen berechnet wird, sondern aus der Verteilung der tatsächlichen Blockzeiten der Vorjahresperiode derselben Saison/Zeitscheibe abgeleitet wird. Häufig wird dafür ein oberes Quantil (z. B. um p80) herangezogen und anschließend auf übliche Zeitraster (z.B. 5 oder 10 Minuten) gerundet. Die konkrete Quantilwahl richtet sich nach Pünktlichkeitszielen, Strecken-/Airportspezifika und gewünschtem Puffer. Bei erhöhten Störfaktoren (z.B. Peak-Taxizeiten, enge Slot-Fenster, wetteranfällige Hubs) tendiert die Quantilwahl nach oben, um Pünktlichkeitsziele zu sichern; auf stabilen City-Pairs bleibt sie näher an p80. [15] [16] Die planmäßige Flugzeit ist dann diese berechnete Blockzeit abzüglich der darin enthaltenen angenommenen Standardrollzeiten.

Die schon optisch erratische Streuung von Off-Block-Zeiten (*Abbildung 4*) und Blockzeiten (*Abbildung 5*) führt im Regelfall zu einer Anflugsequenz, die nicht beliebig vorab gestaltet werden kann, sondern sich hauptsächlich aus den Einflogzeiten in die TMA und aus der dann folgenden Restflugstrecke, Speedvorgaben und den Eindrehpunkten auf den Endanflug ergibt.

Dabei hat die Reihenfolge der auf dem Endanflug befindlichen Flugzeuge einen erheblichen Einfluss auf die individuelle Abstandhaltung und in Summe auch auf die erreichbare Landerate.

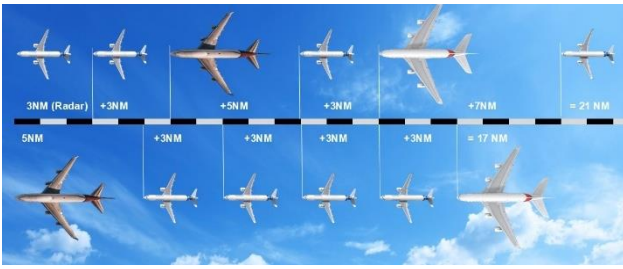


Abbildung 6: Zwei Beispiele für einen identischen Flugzeugmix im Anflug auf eine Piste, in zwei verschiedenen, willkürlichen Anflugreihenfolgen

Abbildung 6 zeigt exemplarisch die Auswirkung einer veränderten Anflugreihenfolge von 6 Flugzeugen auf eine Piste. Der Unterschied in der summierten Mindeststaffelung (modellhafte Betrachtung) beträgt in diesem Fall 4 NM – das ist in etwa der räumliche Bedarf für eine zusätzliche Landung.

Zur Quantifizierung dieses Effekts der überwiegend willkürlichen Anflugreihenfolge wurde eine permutative Analyse auf Basis typischer Verkehrskonstellationen an einem Hub-Flughafen wie Frankfurt durchgeführt (Abbildung 7). Dabei wurden 14 Flugzeuge mit teilweise unterschiedlichen Wake-Turbulence-Kategorien gemäß International Civil Aviation Authority (ICAO) [17] kombiniert, darunter ein Airbus A380 (Weight Type Category (WTC) SUPER HEAVY), vier Widebodies wie der B777-300ER (WTC HEAVY) und 10 Narrowbodies wie A320 oder B737-800 (WTC MEDIUM).

Für jede Variation der Anflugreihung wurde die resultierende Gesamtlänge der Anflugsequenz in nautischen Meilen berechnet, basierend auf den jeweils notwendigen minimalen Staffelungsabständen zwischen den Flugzeugpaaren gemäß ICAO und gemäß RECAT EU [18].

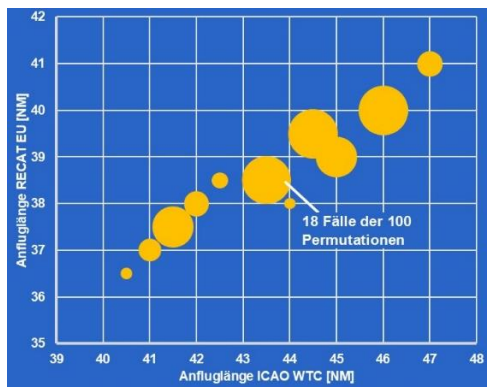


Abbildung 7: Summierte Mindestabstände von 14 Anflügen in unterschiedlichen Anflugfolgen, 100 Permutationen. Zugrunde gelegt für Wirbelschleppenminima wurden zum einen die Abstandswerte gemäß ICAO (X-Achse) und zum anderen die Werte von RECAT EU (Y-Achse)

Die Analyse umfasste insgesamt 100 Permutationen und zeigt eine deutliche Streuung der Gesamtlänge der simulierten Sequenz: Die minimale Anfluglänge bei ICAO-Werten lag bei 40,5 NM, die maximale Länge betrug 47 NM. Der Mittelwert lag bei 44,15 NM bei einer Standardabweichung von 1,62 NM. Die Differenz zwischen ungünstigster und günstigster Sequenz entsprach somit 6,5 NM – eine Distanz, in der je nach Gewichtsklasse ein bis zwei zusätzliche Landungen stattfinden könnten. Bei einem ursprünglich von der University of Westminster berechnetem Durchschnittswert von 130 EUR pro Minute Verspätungskosten [19] [20] entspricht dieser Spread von rund 2,5 Minuten einem Unterschied von bis zu 325 EUR allein für den letzten der 14 Anflüge, ausschließlich verursacht durch die Reihenfolgeeffekte.

Auf Basis von Wirbelschleppenabständen basierend auf RECAT EU ist die Anfluglänge der 14 Arrivals deutlich kürzer (Mittelwert 39 NM) und liegt zwischen 36,5 und 41 NM. Die Standardabweichung liegt deutlich niedriger als unter ICAO-Regime bei 1,03 NM. Bezogen auf Verspätungskosten bedeutet dieser maximale Spread bei RECAT-EU-Abständen rund 244 EUR Mehrkosten für den letzten der 14 Anflüge.

Die in der Permutationsanalyse ermittelten Spannweiten wirken nicht isoliert: Unter Bedingungen hohen Verkehrsaufkommens mit kontinuierlichem Nachschub an Anflügen addiert sich der Effekt fortlaufend zur virtuellen Warteschlange in der TMA. Während sich bei moderater Nachfrage einzelne Sequenzineffizienzen noch abpuffern lassen, führt eine solche Verschiebung bei annähernder Sättigung der Runway-Kapazität zu einem linearen Transfer in zusätzliche Delay-Minuten für nachfolgende Flüge. Im Extremfall können aus einem initialen Nachteil von 2,5 Minuten für den 14. Flug eines Clusters über eine Stunde Betrieb kumulative Verspätungen im zweistelligen Minutenbereich entstehen. Damit wird sichtbar, dass ineffiziente Sequenzierung im Sinne der Warteschlangentheorie nicht nur Einzelflüge trifft, sondern durch emergente Verstärkungseffekte systemweite Auswirkungen auf die Stabilität des Endanflugstroms hat.

Die Ergebnisse unterstreichen das Potential, Anflugsequenzen auf ein mögliches Abstandsminimum der Sequenz zu optimieren, was eine geeignete Systemunterstützung voraussetzt. Automatisierte Sequencing Tools könnten – ergänzt um probabilistische Bewertungsmetriken – eine resilientere Anflugsteuerung ermöglichen, die selbst unter nicht vorhersehbaren Ereignissen (z. B. Go-Around, Mixed Mode Betrieb, Schlechtwetterkorridore) eine möglichst plankapazitätsnahe Abwicklung gewährleistet.

Auf der anderen Seite illustriert die erratische Verteilung von Anflügen und damit verbundenen Mindestabständen, dass es keine deterministisch planbare oder langfristig prognostizierbare Kapazität einer Piste gibt, sondern diese diversen, teilweise zufälligen Einflüssen unterworfen ist, die oft jenseits der Kompensierbarkeit liegen und teilweise schon Stunden zuvor ihren Ausgang genommen haben.

3.2 Sicherheitsreserven im Mixed-Mode-Betrieb: Streuung in der Umsetzung von Startfreigaben

Im Mixed-Mode-Betrieb – also dem abwechselnden Betrieb von Starts und Landungen auf einer gemeinsamen Piste – sind präzise zeitliche Abläufe von zentraler Bedeutung für Sicherheit und Kapazität. Bereits geringfügige Abweichungen in der Umsetzung von Startfreigaben können dabei zu Pistenkonflikten oder unnötigen Verzögerungen für den Anflugverkehr führen. Ein kritischer Punkt ist die Phase unmittelbar nach Erteilung der Startfreigabe, in der das Flugzeug den Startlauf einleitet: Sie hängt maßgeblich von der Reaktion der Cockpitbesatzung ab, der Schubeinstellung, der Flugzeugmasse, atmosphärischen Bedingungen sowie dem Pistenzustand.

Landefreigaben dürfen erst erteilt werden, wenn die Bahn frei ist (der vorhergehende Startvorgang ist abgeschlossen und die Piste frei geräumt). Die zeitliche Unsicherheit der Startumsetzung (Reaktion + Startlauf + Runway Occupancy Time (ROT)) wirkt damit direkt auf die Freigabefähigkeit für den nachfolgenden Anflug. Das unterstreicht die Notwendigkeit quantilbasierter Puffer anstelle fixer Punktzeiten.

Untersuchungen [21] [22] zeigen, dass die Reaktionszeiten auf ATC-Freigaben in dieser Phase nicht konstant sind, sondern deutliche Streuungen aufweisen. In der Praxis bewegen sich Mittelwerte typischerweise zwischen 4 und 6 Sekunden, wobei jedoch Einzelfälle mit Verzögerungen von 10 Sekunden und mehr beobachtet wurden. Auch groß angelegte europäische Analysen der European Organisation for the Safety of Air Navigation (EUROCONTROL) [23] bestätigen dieses Bild: An großen europäischen Flughäfen beträgt die durchschnittliche Zeitspanne von der Takeoff-Clearance bis zum Beginn der Rollbewegung etwa 11 Sekunden, während optimierte Verfahren einen Wert von rund 7 Sekunden als realistisch und kapazitätssteigernd ansehen - wenn eine Einhaltung (ggf. auf technischer Basis) garantiert ist. Eine Absenkung in diesen Bereich kann rechnerisch bis zu zwei zusätzliche Abflüge pro Stunde ermöglichen. Andere Studien [22] weisen zudem darauf hin, dass Unterschiede zwischen Narrowbody- und Widebody-

Flugzeugen oder technische Faktoren wie Triebwerk-Spool-Up-Zeiten zusätzliche Varianz erzeugen. Auch Lindner [24] kommt in einer exemplarischen Analyse von DFS-Radardaten zu dem Ergebnis, dass die mittlere Reaktionszeit der Cockpitbesatzungen bei etwa sechs Sekunden liegt.

Wichtig ist jedoch, dass für die sicherheitsrelevante Abstandshaltung gegenüber einer nachfolgenden Landung nicht allein die Reaktionszeit zählt. Eine Landefreigabe darf erst dann erteilt werden, wenn das startende Flugzeug nicht nur den Startlauf begonnen, sondern tatsächlich abgehoben und das Bahnende überflogen hat.

Während sich der Einfluss des Flugzeugmusters (z. B. A320 versus B777) und der Triebwerkscharakteristik prinzipiell aus dem (zuweilen fehlerbehafteten) Flugplan und der Einfluss des individuellen Flugzeugs aus der MODE-S-Adresse ableiten lässt, bleiben andere Faktoren für die Flugsicherung vage bis unbekannt: z.B. der aktuelle Beladungszustand oder die Wahl einer reduzierten Startschub-Einstellung („Flex Thrust“). Gerade diese Parameter beeinflussen maßgeblich die Länge des Startlaufs – und sind ex ante nur der Cockpitcrew bekannt, da die Parameter dafür oft erst im Taxi-Out gesetzt werden.

Studien wie „Modelling Flexible Thrust Performance for Trajectory Prediction Applications in ATM“ (ATM = Air Traffic Management) zeigen, dass man diese Effekte mit Modellen auf Basis der angenommenen Temperatur, dem Gewicht und der Außentemperatur recht gut abschätzen kann. [25]

Empirische Auswertungen der ROT zeigen, dass vom Beginn des Startlaufs bis zur vollständigen Freigabe der Piste erhebliche Unterschiede bestehen. Für Narrowbody-Jets wie A320 oder B737 liegen Mittelwerte typischerweise zwischen 35 und 45 Sekunden, während Widebody-Muster wie B777 oder A350 im Bereich von 50 bis 70 Sekunden und mehr liegen. [26]

Muster	Mittelwert ROT (s)	5–95 % Perzentil (s)
A320	38–42	30–50
B737-800	36–40	28–48
A330-200	50–60	42–70
B777-300ER	55–65	45–75
A380	65–75	55–85

Tabelle 1: Exemplarische Departure ROT nach Flugzeugmuster [26]

Damit wird deutlich: Die für den Mixed-Mode entscheidende Unsicherheit liegt weniger in der

punktuellen Reaktion des Piloten, sondern in der gesamten, zufallsbehafteten Zeit bis zur Bahnfreigabe. Für die Flugsicherung bleibt dies eine stochastische Größe, deren Variabilität in die Pufferbemessung eingehen muss.

Die zentrale Frage lautet daher: Wie groß muss der zeitliche Staffelpuffer im Mixed Mode bemessen sein, um eine gegebene Wahrscheinlichkeit für konfliktfreie Operationen zu erreichen? Und: Welches Risiko verbleibt bei einer festgelegten Pufferzeit, wenn die Reaktionszeiten und ROT einer stochastischen Verteilung folgen?

Eine Modellierung dieser Situation erfordert die Abbildung der Response Time und ROT als Zufallsvariablen. Ausgehend von einem geplanten Zeitintervall zwischen Start und Landung lässt sich berechnen, wie hoch die Wahrscheinlichkeit ist, dass der Startvorgang rechtzeitig abgeschlossen ist. Je nach Zielwert für die akzeptierte Restrisikoquote (z. B. 99,9 %) ergibt sich daraus ein erforderlicher Puffer – der wiederum die sicher nutzbare Kapazität der Bahn senkt.

Im Folgenden wird die Wahrscheinlichkeitsverteilung der Summe Reaktionszeit + ROT für einen Airbus A320 berechnet, basierend auf einem Monte-Carlo-Modell mit 100.000 Läufen, *Abbildung 8* zeigt das Ergebnis als Grafik.

Piste	Flughafen Frankfurt, Sommer, Runway 25 L
Reaktionszeit (Clearance → Rollbeginn)	Für Frankfurt (stark frequentierter Hub, eingespielte Crews) ist realistisch: Median ca. 7 s, Mittelwert ca. 9–10 s.
ROT (Rollbeginn → Bahn frei)	Mittel ca. 41–44 s, 5–95 % ≈ 32–55 s. Im Sommer (hohe Temperaturen → geringere Triebwerksleistung, längerer Startlauf): leichte Verschiebung nach oben. Annahme: Median ca. 42 s, Mittel ca. 44 s, Streuung ±12 s.
Summe (Reaktion + ROT)	Erwartungswert: ca. 52–55 s , 95 %-Quantil: ca. 70 s , Ausreißer können in Einzelfällen bis knapp 90 s auftreten.

Tabelle 2: Annahmen für die Modellierung der sicheren Abstandshaltung zwischen Arrival und Departure

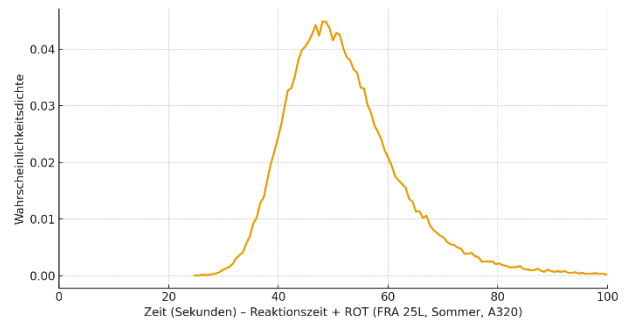


Abbildung 8: Verteilung von Reaktionszeit + ROT (x bis 100 s) für einen Airbus A320, Flughafen Frankfurt, Runway 25L, Sommerbetrieb

G/A-Wahrscheinlichkeit	Nötiger zeitlicher Abstand
bis 0,1%	125,7 s
bis 1%	89,1 s
bis 5%	71,8 s

Tabelle 3: Referenz-Quantile der simulierten Summenverteilung (Reaktion + ROT)

Die ermittelte zeitliche Verteilung für unterschiedliche G/A-Wahrscheinlichkeiten lässt sich unter Annahmen für eine (konstante) Endanfluggeschwindigkeit des folgenden Anflugs in die minimale Distanz zur Schwelle zum Zeitpunkt der Startfreigabe umrechnen.

Die nachfolgende *Tabelle 4* zeigt den minimal erforderlichen Abstand des anfliegenden Flugzeugs zur Landeschwelle zum Zeitpunkt der Startfreigabe für einen wartenden A320 auf FRA 25L (Sommerbedingungen). Die Abstände ergeben sich in Abhängigkeit von der Anfluggeschwindigkeit und einer vorgegebenen maximalen Wahrscheinlichkeit für einen Go-Around.

	Endanfluggeschwindigkeit (GS) in KT		
	120	130	140
G/A-Wahrscheinlichkeit max:	Nötiger Abstand zur Landeschwelle in NM bei Startfreigabe:		
0%	10,9	11,8	12,8
0,1%	4,2	4,5	4,9
1%	3,0	3,2	3,5
5%	2,4	2,6	2,8

Tabelle 4: Nötiger Abstand für definierte G/A-Wahrscheinlichkeiten in Abhängigkeit der Endanfluggeschwindigkeit des folgenden, landenden Flugzeugs

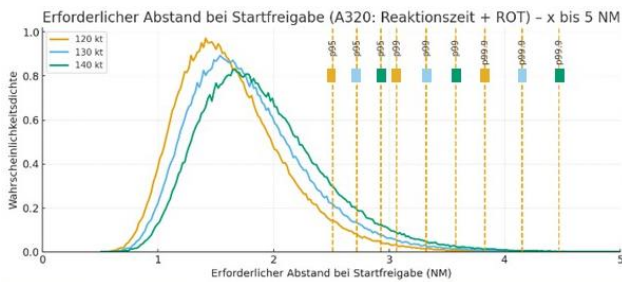


Abbildung 9: Verteilung notwendiger Abstände mit Quantilen für 0,1%, 1%, 5% G/A-Risiko, Airbus A320, Frankfurt Runway (RWY) 25L, Sommerbetrieb¹

In der Grafik (Abbildung 9) wird ersichtlich, dass allein eine um 20 KT verringerte Anfluggeschwindigkeit über Grund (GS) des landenden Fluges den erforderlichen Abstand um rund 0,5 NM oder etwa 15% reduziert. Diese geringe GS kann ihren Grund in einer z.B. aufgrund geringer Masse vergleichsweise um 20 KT geringen Indicated Air Speed (IAS) haben – oder aber z.B. in einer Gegenwindkomponente von 20 KT.

Auswirkungen auf Automatisierung im Bereich Endanflugassistentz

Klassische Automatisierung operiert meist deterministisch. Sie optimiert Zeitfenster, berechnet Zielstarts und synchronisiert mit geplanten Anflügen. Doch ohne Einbeziehung realer Reaktionsstreuungen bleibt ein zentrales Risiko unbeachtet. Die Annahme „das Flugzeug startet genau zum geplanten Zeitpunkt“ ist im Mixed Mode keine verlässliche Grundlage – und kann zu unerwünschten Situationen führen, wenn ein Flugzeug die Piste verspätet verlässt und der nächste Anflug bereits nahe der Schwelle ist. Auf der anderen Seite wäre es nachteilig, regelmäßig sehr bzw. zu große Puffer für große Reaktionszeiten und ROT einzuplanen, da damit „Betonkapazität“ verschenkt wird.

Ein echter Fortschritt wäre eine unsicherheitsrobuste Automatisierung, die nicht nur Zeitpunkte, sondern Wahrscheinlichkeitsräume berücksichtigt – etwa durch:

- **prädiktive Modelle** für Response-Zeitverhalten einzelner Cockpits (z. B. mit Künstlicher Intelligenz (KI) auf Basis historischer Performance),

- **dynamische Sicherheitsmargen** je nach Flugzeugtyp, Crew-Erfahrung oder Wetterlage,
- **automatisierte Monitoring- und Decision-Support-Systeme**, die bei Abweichungen adaptive Anpassungen vorschlagen.

In der ARIS-Strategie des Single European Sky ATM Research (SESAR) JU [13] wird genau diese Notwendigkeit adressiert: Zukunftsfähige Systeme sollen „resilient under operational variability“ und „trustworthy under uncertainty“ sein. Das bedeutet nicht, dass der Start selbst automatisiert werden muss – sondern dass die Planung und das operative Regelwerk künftig probabilistisch-informiert reagieren muss. Ein statischer 90-Sekunden-Puffer ohne Rücksicht auf Variabilität reicht in solchen Umgebungen nicht mehr aus.

Automatisierung kann also helfen – nicht durch Verdrängung menschlicher Entscheidung, sondern durch die systematische Integration von Unsicherheit in das Entscheidungsmodell.

Eine sinnvolle Ergänzung einer systemgestützten Assistenz könnte dabei die sensorische Erweiterung für das Lotsenpersonal sein. Derzeit ist der Beginn des Startlaufs für Fluglotsen schwer zu erkennen – zumal in vielen Fällen (z.B. in Frankfurt bei Betriebsrichtung West) der Tower buchstäblich meilenweit von der Bahnschwelle entfernt ist. Auch auf der Darstellung der Bodenlage ist der Beginn des Startlaufs schwer abzuschätzen. Eine vergleichsweise kostengünstige Lösung könnte darin bestehen, den Beginn des Spool-Up mit Mikrofonen am Bahnkopf zu erfassen und den Lotsen in der Logik einer „Ampel“ anzuzeigen. Diese Information wäre auch relevant für das Training passender KI-Algorithmen.

Mit der gezielten Auswertung zusätzlicher Sensorik und Daten (Echtzeitinformationen zu Masse, Prognose von Schubeinstellungen und Wind oder die Nutzung bestehender Datenquellen wie MODE-S zur Abschätzung der Triebwerksperformance) lassen sich Varianzen deutlich verringern. Anstelle breiter Verteilungen mit schwer planbaren Ausreißern entstehen schmalere, besser prognostizierbare Kurven. Damit werden Quantile wie p95 belastbarer, was eine robustere Staffelungsbemessung bei gleichbleibender Sicherheit ermöglicht.

¹ Die im Vergleich zur Zeitskala höhere Dichte bei der Darstellung über der Distanz erklärt sich aus der unterschiedlichen Skalierung: Da die Wahrscheinlichkeitsdichte stets so normiert ist, dass die Fläche unter der Kurve den Wert 1

ergibt, führt die Transformation von Sekunden auf nautische Meilen (mit engerem Wertebereich) zwangsläufig zu höheren Maximalwerten der Dichte.

3.3 Einfluss von Kommunikationsblocking auf die Steuerbarkeit im Anflug

In hochbelasteten Anflugsektoren ist die effektive Steuerung von Verkehrssequenzen im Downwind-Bereich in Richtung Intermediate und Final Approach wesentlich für die Maximierung der Landerate und die Aufrechterhaltung stabiler Betriebsbedingungen. Besonders auf der Position des „Director“ (zuständig für den Endanflug inkl. Eindrehmanöver) erfordert die Gestaltung effizienter Sequenzen eine fortlaufende Intervention durch Kursfreigaben, Speed Adjustments und/oder Vorstaffelungsanpassungen.

Diese Interventionen sind jedoch auf einen knappen Kommunikationskanal angewiesen – in der Regel einen einzelnen Sprechfunkkanal pro Frequenz. Bereits temporäre Belegungen des Kanals führen zu so genannten Blocking Events, die die unmittelbare Steuerbarkeit des verantworteten Verkehrs einschränken.

Ein besonders kritischer Fall tritt auf, wenn Initial Calls – also Erstmeldungen neu in den Sektor eintretender Luftfahrzeuge – die Frequenz blockieren. Diese Calls werden von den Cockpitbesatzungen autonom bei Einflug vorgenommen, sind häufig länger als routinierte Anweisungen, beinhalten Zusatzinformationen (z. B. Wetter, Transponderstatus, Standby Calls, WTC) und können durch Nachfragen verlängert werden.

Fällt ein solcher Call in eine zeitkritische Phase (etwa kurz vor dem Base Turn eines Fluges auf dem Downwind), kann das optimale Zeitfenster für eine notwendige Freigabe verloren gehen. In der Praxis zeigt sich dies in verlängerten Downwinds oder sogar Go-Arounds, wenn Staffellungen nicht mehr rechtzeitig herstellbar sind.

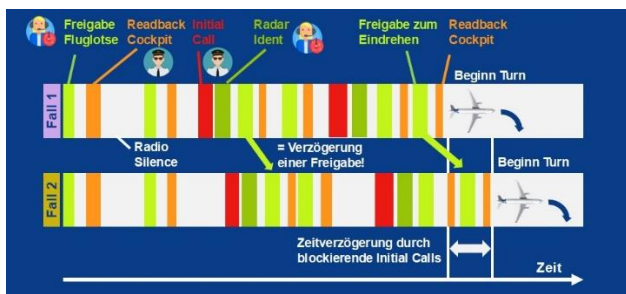


Abbildung 10: Zeitstrahlen veranschaulichen den Effekt blockierender Initial Calls: Während im ungestörten Ablauf (Fall 1) Freigaben und Readbacks einander ohne Störung folgen, verdrängen die Initial Calls bei ungünstiger Zeitenlage (Fall 2) reguläre Freigaben und führen zu Verzögerungen bei der Manövereinleitung.

Modellierung als Single-Server-Warteschlange

Dieser beschriebene Kommunikationskanal wurde in warteschlangentheoretischer Analogie als Single-Server-Queue modelliert: Ankunftsprozess der Radiotelefonie-(R/T)-Ereignisse als Poisson-Prozess ($\approx 9,2$ Ereignisse/Minute, davon $\approx 13\%$ Initial Calls), Servicezeiten triangular verteilt (Initial Calls 3–7 s, sonstige 2,5–6,5 s), Bedienpolitik First-Come-First-Served (FCFS), nicht-präemptiv.

Geplante Freigaben wurden deterministisch im Mittel alle 15 s simuliert. Verzögerungen treten dann auf, wenn der geplante Zeitpunkt einer Freigabe in eine laufende Belegung durch einen Readback oder einen Initial Call fällt – dann verschiebt sich der Beginn der Freigabe bis zum Ende des laufenden Calls.

Methodische Annahmen und Gültigkeitsbereich

Die für die Modellierung und Simulation verwendeten Parameter repräsentieren stark ausgelastete Anflug-/Endanflug-Sektoren an großen Hubs (Tagesrand- und Peak-Phasen). Für die Funk- und Freigabeprozesse wird eine mittlere R/T-Ereignisrate im oberen Lastbereich angesetzt; Erstmeldungen („Initial Calls“) weisen längere, teils variabelere Sprechzeiten auf als Standardanweisungen. Downwind-Geschwindigkeiten wurden typisch auf ≈ 210 knots (KT) modelliert; für Mixed-Mode-Betrachtungen wurden muster- und saisonabhängige Reaktionszeiten und ROT-Verteilungen genutzt.

Die Ergebnisse gelten damit primär für hochbelastete Betriebsfenster; in Off-Peak-Phasen sind die Effekte geringer, die qualitative Logik (Quantile statt Punktwerte, Risikobudgets entlang der Kette *Wahrnehmen* → *Übertragen* → *Umsetzen*) bleibt bestehen.

Simulationsergebnisse

Nach 200 Replikationen von jeweils 60 Minuten (≈ 48.000 geplante Freigaben) ergab sich folgendes Bild:

Kennwert	Wert
Anteil verzögerter Freigaben	$\approx 43\%$
Median-Verzögerung (bezogen auf alle Freigaben)	0,15 s
p90	3,86 s
p95	4,45 s
p99	5,34 s
Max beobachtet	6,63 s
Initial-Call-Paare pro Stunde	≈ 46

Blockdauer Initial-Paar (Mittel)	≈ 8,2 s
Blockdauer Initial-Paar (p95)	≈ 10,7 s
Anteil Freigaben im Doppel-Block	≈ 8 %

Tabelle 5: Simulationsergebnisse über die Auswirkungen von Initial Calls und damit verbundener Blockierung der Kommunikation für Freigaben

Die Verteilung (Abbildung 11) der Verzögerungen weist einen markanten Peak auf: Rund 57% aller geplanten Freigaben erfolgen ohne jede Verzögerung (Peak bei 0 s). In den übrigen ≈ 43% tritt dagegen eine Verzögerung auf, die gleichmäßig zwischen etwa 0,1 s und 6 s verteilt ist.

Ein zweiter Schwerpunkt ist dabei nicht erkennbar. Dieses Muster verdeutlicht den systemischen Charakter der Single-Server-Logik: Entweder der Kanal ist sofort frei, oder es entsteht eine – im Einzelfall zufällige – Wartezeit bis zur nächsten Gelegenheit. Dadurch können Blocking-Events die Planbarkeit abrupt und unverhältnismäßig stören.

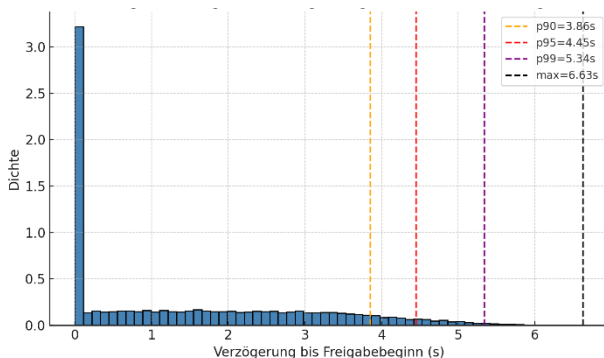


Abbildung 11: Histogramm der Verzögerungen bei Freigaben (200×60 min, high load).

Implikationen

Die Gesamtkapazität der Runway wird nicht allein durch ihre physikalisch-technischen Parameter bestimmt, sondern auch durch die kommunikative Verarbeitungskapazität. Klassische deterministische Kapazitätsmodelle unterschätzen diesen Faktor, da sie von permanenter Verfügbarkeit der Frequenz ausgehen. Erst die probabilistische Modellierung zeigt, dass Kommunikationsblocking eine signifikante, stochastische Restriktion der Steuerbarkeit darstellt. Besonders der Sprung zwischen ungestörten und geblockten Freigaben illustriert die Abruptheit, mit der Planbarkeit verloren gehen kann.

Für die operative Steuerung auf dem Downwind bedeutet dies: Bei einer typischen Geschwindigkeitsführung von etwa 210 KT kann eine verspätete Freigabe

zum Turn auf Base/Final mit einer Wahrscheinlichkeit von ~43 % eintreten. Die dadurch zusätzlich geflogene Strecke liegt im p90 bereits ≈ 0,143 Nautische Meile (NM), im p95 bei ≈ 0,177 NM und im p99 bei ≈ 0,248 NM. In seltenen Spitzenfällen werden ≈ 0,387 NM an potenziellem Abstand zur vorlaufenden Maschine verschenkt, sofern der Lotse keine Gegenmaßnahmen zur Kompensation einleitet.

Diese Verluste verringern die Freiheitsgrade für den zeitkritischen Base-Turn und engen die Möglichkeiten zur Sequenzoptimierung ein. Praktisch resultiert daraus eine erhöhte Wahrscheinlichkeit, dass Staffeln erst später – z.B. mit späteren und deutlicheren Speed-Reduktionen – wiederhergestellt werden können; im Engpassbetrieb steigt damit indirekt auch das Risiko für staffelungsbedingte Go-Arounds.

Im Flugsicherungsbetrieb überlagert sich der Effekt durch Funkblockaden mit den vorab in diesem Artikel bereits dargestellten Effekten einer verzögerten Reaktion auf Freigaben - *Abbildung 12* zeigt die berechnete Verteilung der Wahrscheinlichkeitsdichte unter Zugrundelegung der o.g. typischen Verteilung von Reaktionszeiten.

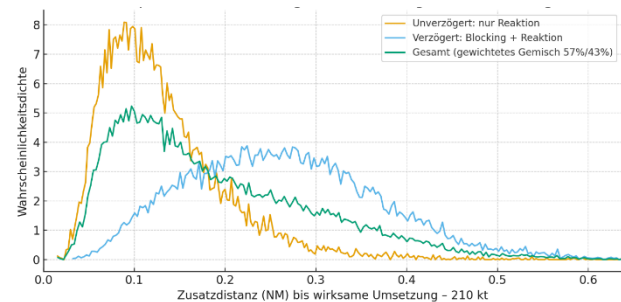


Abbildung 12: Zusatzdistanz bis zur wirksamen Umsetzung von avisierten Freigaben zum Eindrehen (Reaktionszeit wirkt bei „blockierten“ und nicht „blockierten“ Freigaben.)

Kapazitätswirkung und ökonomischer Gegenwert

Verzögerungen bei Eindrehfreigaben wirken sich nicht nur auf einzelne Flugwege aus, sondern auch auf die gesamte Landerate – und zwar dann, wenn die Verzögerungen Anflüge betreffen, die sich in der kritischen Phase der Sequenz befinden, also unmittelbar auf den Headway der nachfolgenden Maschine wirken. Für die kritischen Anflüge beträgt die mittlere Zusatzverzögerung rund 3 s. Da etwa 70 % aller Anflüge in diese Gruppe fallen (Festlegung für die Modellierung), ergibt sich für die Gesamtheit der Anflüge eine gewichtete Zusatzverzögerung von 2,3 s. Der mittlere Headway steigt damit von 90 s auf 92,3 s, was rechnerisch einer Landerate von rund 39 Landungen pro

Stunde entspricht – also etwa eine Landung pro Stunde weniger als im ungestörten Fall.

Kennzahl	Wert	Einheit / Bedeutung
Zeitkritischer Anteil g	70%	Anteil der Anflüge, die ohne zusätzlichen Puffer geführt werden
Zusätzliche Zeit pro Anflug $E[X]$	3,24 s	Mittlere Zusatzverzögerung je kritischem Anflug
Varianz $Var[X]$	3,51 s ²	Streuung der Verzögerungszeiten
Neuer mittlerer Headway $E[T]$	92,3 s	Durchschnittlicher Abstand zwischen Landungen
Erwartete Landerate $E[N]$	39,0	Landungen pro Stunde – theoretischer Mittelwert bei Berücksichtigung aller Verzögerungen
Standardabweichung $SD[N]$	0,089	Landungen pro Stunde – typische Schwankungsbreite der Landerate in einem 1-h-Fenster
95%-Intervall (1 h)	38,8 – 39,2	Landungen pro Stunde – Bereich, in dem die beobachtete Landerate mit 95% Wahrscheinlichkeit liegt (kann leicht unter oder über dem Mittelwert liegen)

Tabelle 6: Auswirkungen von Reaktionszeiten und Kommunikationsblockierungen auf die rechnerische Landerate einer Piste

Ökonomischer Gegenwert einer Landung pro Stunde

Zwei in *Tabelle 6* illustrierte Perspektiven sind hilfreich, um den Wert einer marginalen Landung/h zu beziffern – unter Zugrundelegung des Kostenansatzes der University of Westminster von ~130 €/Min Verspätung. [19]

Kurzfristige Marginalbetrachtung (Headway-Zuwachs)	Der erwartete Zusatz-Headway $g \cdot E[X] \approx 2.27$ s führt bei ≈ 39 Landungen/h zu einer kumulierten zusätzlichen Flugzeit von ~1,48 Minuten je Stunde (≈ 192 € pro Stunde für alle Anflüge zusammen).
Anhaltende Übernachtfrage (Spillover-Sicht)	Bei einer Nachfrage von 40/h und einer effektiven Kapazität von ≈ 39 /h entsteht ein systematischer Rückstand. In der ersten Peak-Stunde entspricht dies rechnerisch einer kumulierten Zusatzverspätung von etwa 30 Minuten im System.

	Hält die Übernachtfrage jedoch länger an, wächst die Warteschlange linear weiter an – in diesem instabilen Dauerzustand entspricht der Rückstand ungefähr einem Flug pro Stunde, d. h. rund 60 Verspätungsminuten. Bewertet mit branchenüblichen Kostenansätzen (≈ 130 €/min) ergibt dies rund 3.800 €/h für eine einzelne Peak-Stunde und bis zu 7.700 €/h bei längerer Übernachtfrage, die allerdings für einen (Hub-)Flughafen untypisch ist.
--	---

Tabelle 7: Konsequenzen einer kurzfristigen und langfristigen Auswirkung der Landeratenreduktion

In transienten Peaks ist der unmittelbare Mehrzeit-Effekt moderat, in persistenten Hochlastphasen dominiert der Spillover-Mechanismus: die fehlende Landung je Stunde akkumuliert sich und erzeugt schnell vierstellige Eurobeträge pro Stunde an systemischer Verspätung. Der „Gegenwert“ einer potentiellen zusätzlichen Landung pro Stunde liegt damit – je nach Betriebszustand – zwischen hunderten Euro (kurzfristig) und mehreren Tausend Euro pro Peak-Stunde (persistente Übernachtfrage).

Konsequenzen für eine weitergehende Automatisierung

Automatisierung muss Kommunikationsverfügbarkeit explizit als stochastisch schwankenden Ressourcenfaktor berücksichtigen. Langfristig ist eine unsicherheitsrobuste Automatisierung erforderlich, die Kommunikationslatenzen und Blocking probabilistisch modelliert und in die Entscheidungslogik integriert – im Sinne des SESAR ATM Master Plan 2025 („air-ground integration under variable latency“).

Off-Nominal-Ereignisse (z. B. Medicals, technische Rückläufe, Sicherheitsereignisse) wirken als exogene Schocks in dieselbe Logik: Sie erhöhen die Varianz, ziehen Gates nach vorn, steigern kurzfristig die Funklast und machen quantilbasierte Puffer statt Punktzeiten erforderlich. Das Framework aus Gates, Quantilen und Risikobudgets bleibt gültig – lediglich die Zielquantile werden für die Dauer des Ereignisses konservativer gewählt.

Entscheidend ist somit nicht die Eliminierung aller Variabilität, sondern ihre bewusste Integration in automatisierte Entscheidungslogiken. Die Zukunft der

Flugsicherung liegt nicht in der Illusion deterministischer Präzision, sondern in probabilistischen Verfahren, die Unsicherheit sichtbar machen und gezielt beherrschbar halten.

3.4 Regelkreise zwischen Mensch und Maschine – Bord und Boden

Flugverkehrskontrolle bedingt eine enge Kooperation zwischen Fluglotsen und Cockpitcrews – gestützt auf kognitive Ressourcen, Systemunterstützung und bilaterale Kommunikation. In der Human-Factors-Forschung existieren komplementäre Rahmenwerke, die diese Abläufe strukturieren: OODA (Observe–Orient–Decide–Act) beschreibt den iterativen Entscheidungszyklus, Endsleys Situation Awareness differenziert *Wahrnehmen–Verstehen–Antizipieren* als Voraussetzung wirksamer Eingriffe, und Rasmussen gliedert menschliches Handeln in skill-, rule- und knowledge-basierte Ebenen. [27] [28]

Kontrolltheoretische Perspektiven klassifizieren „Mensch-im-Regelkreis“-Architekturen nach Kopplung und Bandbreite. Gemeinsam erklären diese Ansätze, warum Latenzen und Variabilitäten auf mehreren Ebenen entstehen – und wie sie die Regelgüte beeinflussen.



Abbildung 13: Regelkreise Orchestrierung Flugverkehr, Kommunikation Fluglotse-Cockpit und Flugführung

Für das Flugverkehrsmanagement lassen sich drei hierarchisch verschachtelte Regelkreise identifizieren (Abbildung 13), die jeweils von deterministischen Abläufen und stochastischen Streuungen geprägt sind:

1. Kognitiver Regelkreis des Fluglotsen

- **Regelgröße:** aktuelles Verkehrslagebild („Picture“), inkl. prognostizierter Trajektorien
- **Regler:** Fluglotse (unterstützt durch ATS-Systeme und Assistenzsysteme)
- **Stellgröße:** geplante Anweisung (z. B. Heading, Speed)

- **Störgrößen:** Bildaufbauverzögerung, Systemverzögerungen, kognitive Einschätzungsunsicherheit

2. Kommunikativer Regelkreis Fluglotse – Cockpit

- **Regelgröße:** wahrgenommene Anweisung
- **Regler:** Cockpitcrew (Verständnis, Rückmeldung, Umsetzung)
- **Regelstrecke:** Sprechfunk mit Antwortzeit
- **Störgrößen:** Frequenzüberlagerung, Missverständnisse, Delay durch Workload und nötige Abstimmungen

3. Dynamischer Regelkreis Flugzeugführung

- **Regelgröße:** Steuerkurs / Geschwindigkeit
- **Regler:** Bordcomputer / Pilot
- **Regelstrecke:** Flugzeugdynamik (z. B. Turn Rate, Speed Change Lag)
- **Störgrößen:** Flugverhalten, Trägheit, meteorologische Einflüsse

Diese drei Kreise wirken – mit jeweils eigener Dynamik – vom Fluglotsen gleichzeitig auf mehrere Flugzeuge, die sich in einer zusammenhängenden, ggf. konvergierenden Verkehrssituation befinden können.

Bereits kleine Verzögerungen in einem der Kreise können den Handlungsspielraum für sichere Flugführung und Abstandshaltung deutlich einschränken. Im Worst Case ergibt sich ein „Verkettungsproblem“: Ein leicht verzögertes Lagebild, eine etwas verspätete Funkauslösung und eine träge Kurseinleitung summieren sich zu einer Konfliktverschärfung, obwohl jedes Teilsystem für sich genommen noch im Toleranzbereich operiert hätte.

Stimmigkeit der Regelkreise: Bedingungen, Bruchstellen und Ansatzpunkte

Die drei Regelkreise – *Wahrnehmen/Entscheiden* (Lotse), *Übertragen/Bestätigen* (Funk) und *Verstehen/Umsetzen* (Cockpit/Flugzeug) – sind in der Praxis seriell gekoppelt und wirken zugleich gegenseitig zeitkritisch. Jede Stufe bringt eigene, teils zufallsbehaftete Latenzen mit: Aufbau und Erhalt des „Picture“ sowie Priorisierung im Tower/Approach, Frequenzbelegung und Readback/Hearback-Risiko im Funk, Aufmerksamkeitswechsel, (ggf. durch Wetter beeinflusste) Flight Management System (FMS)-Eingaben und Systemträgheiten im Cockpit.

Diese Latenzen summieren sich und treffen im Endanflug auf eng begrenzte Entscheidungsfenster (z.B. Base-Turn, Landefreigabe vor definierter Gate-Diszanz). So entstehen aus an sich kleinen Verzögerungen schnell Koinzidenzen, die den

Handlungsspielraum abrupt verkleinern (z. B. verspätete Freigabe → späte Umsetzung → Go-Around) und den Lotsen oder die Lotsin verärgern.

Typische Bruchstellen – und warum sie auftreten.

1. **Frequenzspitzen** (*Initial Calls*, Rückfragen, verkettete Freigaben) verschieben den Sprechzeitpunkt hinter das operative Fenster.
2. **Botschaftskomplexität** verlängert Redezeiten und erhöht das Risiko von Hearback/Readback-Fehlern; Korrekturen kosten zusätzliche Kanalzeit.
3. **Medienbrüche** (Voice-Anweisung → manuelle FMS-Eingabe) verzögern die Wirksamkeit und binden Cockpit-Ressourcen im falschen Moment.
4. **Asynchrone Sichten** von Boden und Bord (Zeitziele, Profile) erzeugen späte Korrekturen statt früher, ruhiger Steuerung.
5. **Mixed-Mode-Schnittstelle:** Start-/Landefenster kollidieren mit Kommunikationsspitzen; die Landefreigabe rutscht hinter das Gate.

Wie diese Regelkreise „harmonisieren“ können

Hoch belastbare und zuverlässige Abläufe entstehen, wenn vier Voraussetzungen gleichzeitig erfüllt sind:

Gemeinsame Zeitbasis mit klarer *latest-safe-clearance*-Marke (*Gate*), die für Boden und Bord sichtbar ist. SESAR-Analysen zeigen, dass eine gemeinsame Zeitbasis in AMAN/DMAN-Systemen Voraussetzung für stabile Sequenzen ist, indem Controller und Piloten dieselbe Zeitreferenz als ‚latest-safe-clearance‘ nutzen können. [29]

Niedrige Kommunikationslast bei hoher Eindeutigkeit (kurze, standardisierte Botschaften; eindeutige Parameter/Callsign). Das SESAR-Projekt PJ.16-04 entwickelte eine semantische Ontologie, die ATC-Sprachbefehle standardisiert (Callsign, Command, Value) und damit eindeutige, kurze Übertragungen ermöglicht. [30]

Fehlerrobuste Bestätigung (Readback mit semantischer Prüfung; hör-/systemseitige Mismatch-Erkennung z.B. unter Nutzung von Spracherkennung). Forschungsprojekte wie MALORCA und HAAWALL zeigen, dass Spracherkennungssysteme Readbacks nicht nur transkribieren, sondern semantisch analysieren können. Dadurch lassen sich Mismatch-Situationen zwischen Freigabe und Rücklesung

automatisch erkennen und markieren, sodass Controller fehlerhafte Bestätigungen unmittelbar korrigieren können. [31]

Direkt wirksame Umsetzung im Cockpit (minimierte Medienbrüche, geringe Eingabe-/Moduswechsel-Last, vorhersehbare Dynamik).

Ansatzpunkte für mehr Stimmigkeit (ohne „Mehr Technik um jeden Preis“)

Kommunikation vereinfachen und vorziehen.

Kurz, standardisiert, parameter-erst: Call-Sign → Freigabe → **konkrete** Parameter → Schlusswort. *Sprech-Planpunkt* bewusst **vor** das Gate legen (mit Vorhalt für typische Frequenzwartezeit).

Quasi-automatischer Readback.

Wo verfügbar: digitale Bestätigung mit Syntax-/Parameter-Check als Ergänzung zur Voice-Rücklesung; Mismatch-Flagging für Lotsen.

Direkte Systemanbindung.

Uplink-to-FMS/Automatic Flight Control System (AFCS) (Heading/Speed/Level/Standard Terminal Arrival Route (STAR)-Amendments) reduziert Medienbrüche und beschleunigt die Wirksamkeit der Freigabe.

Geteilte Zeitsicht anstelle von Punktzeiten.

Gate-Anzeige und „nächste freie Sendelücke“ im Lotsentool; auf Cockpit-Seite klar definierte Crew-Gates (Distanz/Zeit/Höhe), ab denen ohne Freigabe GA einzuleiten ist.

Medien bewusst trennen.

Zeitkritisches (Vectors, „clear to land“) Voice-first; *nicht-zeitkritisches* (Level-/Routen-Anpassungen mit Vorlauf) CPDLC (Controller-Pilot Data Link Communication – vulgo „Datalink“) first – mit klarer Fall-back-Policy.

Training auf Engpass-Muster.

Phraseologie-Drills für Peak-Phasen, Gate-SOPs (beide Seiten), bewusste Begrenzung der Botschaftskomplexität.

Essenz

Die Regelkreise werden nicht „stimmig“, indem einzelne Komponenten isoliert präziser werden. Entscheidend ist die End-to-End-Synchronisierung von Zeitfenstern (Gate-Logik), Medien (Voice/Datalink in ihrer Stärke) und Umsetzung (FMS-Anbindung, geringe Eingabelast) – flankiert von einfachen, robusten Verfahren und einer geteilten Zeitsicht. Dort, wo diese Kopplung gelingt, verschwinden Koinzidenzen aus Frequenzblocking, Readback und Umsetzung

weitgehend aus den kritischen Fenstern; wo sie fehlt, bleiben selbst „kleine“ Latenzen betriebsbestimmend.

4. Grenzen deterministischer Systeme und Potenziale probabilistischer Unterstützung

Die operative Flugsicherung ist nicht beliebig präzise planbar: Anflugreihenfolgen verschieben sich, Funkfenster öffnen und schließen sich, Reaktions- und Umsetzungszeiten variieren – und gerade in den entscheidenden Sekunden kumulieren diese Varianzen. Modelle, die von festen Trajektorien oder störungsfreier Kommunikation ausgehen, stoßen hier methodisch an Grenzen. Effizienz und Sicherheit entstehen nicht durch Eliminierung aller Störungen, sondern durch Robustheit gegenüber unvermeidbaren Abweichungen.

Daraus folgt ein Perspektivwechsel: von Punktwerten zu Wahrscheinlichkeitsräumen. Unvermeidbare Zufallsstreuung (*aleatorisch*) ist von reduzierbaren Wissenslücken (*epistemisch*) zu unterscheiden. Für zeitkritische Entscheidungen sind Quantile relevanter als Mittelwerte – die entscheidende Frage lautet: *Mit welcher Wahrscheinlichkeit überschreiten wir ein sicherheitskritisches Gate?* Sicherheitsmargen werden so zu Risikobudgets entlang der Kette *Wahrnehmen – Entscheiden – Übertragen – Umsetzen*.

Damit verbunden ist das Konzept des *calibrated trust*: Vertrauen in Assistenzsysteme entsteht nicht blind, sondern durch kalibrierte Prognosen mit transparentem Gültigkeitsbereich. Der Lotse bleibt Supervisory Decision Maker – unterstützt durch probabilistische Verfahren, die Unsicherheit sichtbar und steuerbar machen.

Der größte Nutzen entsteht dort, wo deterministische Annahmen heute zu knappen Fenstern führen: Sequenzierung und Spacing im Endanflug, Gate-Entscheidungen, Frequenznutzung oder Runway-Zeiten. Praktisch bedeutet das: Sprechzeitpunkte werden so gewählt, dass Freigaben mit hoher Wahrscheinlichkeit vor dem Gate ankommen; Anflugreihenfolgen werden nach Quantilen optimiert; Mixed-Mode-Fenster werden als Konfidenzintervalle statt fixer Werte geplant.

Für die Gestaltung gilt:

- **möglichst einfache Human Machine Interface (HMI)-Darstellungen** (z. B. p50/p90/p95, Ampel für Exzedenzwahrscheinlichkeiten),
- **Reduktion von Medienbrüchen** (FMS-loadable Clearances),

- **Medientrennung nach Zeitkritik** (Voice-first für Zeitkritisches, Datalink für Planbares),
- **gemeinsame Zeitsichten** für Boden und Bord,
- **konservative Degradation bei Datenmangel.**

Mensch und KI – komplementäre Stärken und Bedingungen für Kooperation

Fluglotsen treffen Entscheidungen unter Unsicherheit: variierende Reaktionszeiten, unvollständige Daten und komplexe Dynamiken prägen ihren Alltag. Sie stützen sich dabei auf Intuition, Erfahrung und das bewusste Einbeziehen von Unsicherheit in ihre Entscheidungslogik. Zukünftig wird entscheidend sein, rechnergestützte Assistenzsysteme an diese Prozesse sinnvoll „anzukoppeln“ – und sie dort rechnerisch nachzubilden, wo Mensch und Maschine ähnliche Entscheidungswege beschreiten.

Im Gehirn wirken für wesentliche Lotsenaufgaben verschiedene Areale: [27]

- *Präfrontaler Cortex*: Abwägen von Handlungsoptionen und Wahrscheinlichkeiten
- *Arbeitsgedächtnis*: paralleles Jonglieren mit Szenarien
- *Limbisches System*: emotionale Marker („Gefühl der Richtigkeit“)
- *Basalganglien*: Mustererkennung, Automatisierung wiederkehrender Abläufe
- *Parietalcortex*: Integration räumlich-visueller Informationen (Situational Awareness)

Im Rechner finden sich funktionale Parallelen:

- *Wahrscheinlichkeitsmodelle*: Ranking von Optionen
- *Datenpuffer und Algorithmen*: parallele Szenarierechnungen
- *Heuristiken und Scores*: schnelle Bewertung von Abweichungen
- *Machine Learning*: Mustererkennung in großen Datenmengen
- *Sensorfusion*: Zusammenführen von Verkehrs- und Wetterdaten für ein konsistentes Lagebild

Dennoch werden kognitive Prozesse von Mensch und Maschine auch zukünftig unterschiedliche Merkmale aufweisen:

Mensch (Fluglotse)	Künstliche Intelligenz
Intuition, Erfahrung, Mustererkennung	Mustererkennung durch Training
Kombination aus rationaler Analyse + emotionalen Markern	Statistische/probabilistische Bewertung

Stark im Umgang mit Unsicherheit und Ausnahmen	Stark bei konsistenten, datenreichen Mustern
Gefahr von Vigilanzabfall bei reiner Überwachung	Gefahr von Bias im Trainingsdatensatz

Herausforderungen für gelingende Kooperation:

- **Aufgabengerechte Programmierung:** Systeme müssen probabilistische Logiken abbilden, nicht nur deterministische Regeln.
- **Gegenseitige Einschätzung:** Mensch und Maschine brauchen Transparenz über ihren jeweiligen Arbeitszustand (Workload, Datenqualität).
- **Valide Daten in Echtzeit für Mensch und Maschine:** Verfügbarkeit, Aktualität und Qualität der Informationen sind kritische Erfolgsfaktoren.
- **Schnittstellengestaltung:** Nur wenn Systeme menschliche Stärken aktivieren und nicht verdrängen, entsteht robustes Mensch-Maschine-Teaming.

5. Fazit und Ausblick

Die Befunde zeigen: Operative Flugsicherung ist streuungsgetrieben. Reihenfolge, Frequenzverfügbarkeit, Reaktionszeiten und Flugzeugdynamik variieren in relevanter Größenordnung – deterministische Punktplanungen greifen hier zu kurz. Robust wird das System erst, wenn Unsicherheit sichtbar gemacht und bewusst bewirtschaftet wird: mit Quantilen, Exzedenzwahrscheinlichkeiten und klar verteilten Risikobudgets entlang der Kette *Wahrnehmen – Entscheiden – Übertragen – Umsetzen*.

Wo Boden und Bord eine gemeinsame, vierdimensionale Trajektorie (4D-Sicht) mit Unsicherheitsbändern teilen, sinkt der Bedarf an späten Eingriffen – messbar in stabileren Fenstern, geringerer Funklast und reduzierter Go-Around-Wahrscheinlichkeit bei gleicher Sicherheit. Kommunikation zwischen Boden und Bord muss künftig noch robuster, eindeutiger und zeitkritischer gestaltet werden: Nur durch klare Zeitfenster, standardisierte und möglichst medienbruchfreie Übertragungswege sowie eine Reduktion der Kommunikationslast lassen sich Verzögerungen und Kapazitätsverluste wirksam begrenzen.

Aus diesen Maßnahmen erwächst ein zeitgemäßes Verständnis von *calibrated trust*: Automatisierung ersetzt den Menschen nicht, sondern stärkt ihn als *supervisory decision maker*. Vorausgesetzt sind kalibrierte Modelle mit transparentem Gültigkeitsbereich, einfach lesbare HMI-Signale und klare Rollen zwischen Assistenz und Aufsicht.

Das Ziel ist eine Flugsicherung, die Effizienz und Resilienz nicht trotz, sondern gerade wegen des bewussten Umgangs mit Unsicherheit steigert. Die Zukunft liegt nicht in der Illusion deterministischer Präzision, sondern in probabilistischen Verfahren, die Streuung in Robustheit verwandeln – und den Menschen befähigen, das „Würfeln der Realität“ zu beherrschen.

Erforderlich ist dafür auch die systematische (und aufwändige) Erschließung und Nutzung zusätzlicher Daten- und Informationsquellen: von Echtzeitparametern wie Masse, Schub und Wind über akustische und sensorische Trigger bis hin zur Auswertung vorhandener Datenströme wie MODE-S. Je vollständiger diese Informationen verfügbar und in kalibrierte Modelle integriert werden, desto schmaler werden die Verteilungen, desto verlässlicher die Quantile – und desto robuster die Planungen und Entscheidungen in der Flugverkehrskontrolle.

Abkürzungsverzeichnis und Referenzen

Abkürzung	Englische Langform	Deutsche Bedeutung
4D	Four-Dimensional Trajectory	Vierdimensionale Trajektorie (3D + Zeit)
ACAM	Airport Capacity Assessment Methodology (EUROCONTROL)	Kapazitätsanalyse-Methode von EUROCONTROL
AFCS	Automatic Flight Control System	Automatisches Flugregelungssystem
AMAN	Arrival Management System	Anflugmanagementsystem
AI	Artificial Intelligence	Künstliche Intelligenz
ATC	Air Traffic Control	Flugverkehrskontrolldienst
ATCO	Air Traffic Controller	Fluglotse/Fluglotsin
ATFM	Air Traffic Flow Management	Verkehrsflussmanagement
ATM	Air Traffic Management	Flugverkehrsmanagement
CPDLC	Controller–Pilot Data Link	Datenverbindung Bord–Boden für Flugverkehrsfreigaben

	Communica- tions	
DFS		DFS Deutsche Flugsicherung GmbH
DLH	IATA designator "LH" (Lufthansa)	IATA-Kennung von Lufthansa
DLR		Deutsches Zentrum für Luft- und Raumfahrt
DLRK		Deutscher Luft- und Raumfahrtkongress
DMAN	Departure Management System	Abflugmanagementsystem
EUROCONTROL	European Organisation for the Safety of Air Navigation	Europäische Organisation zur Sicherung der Luftfahrt
FCFS	First Come, First Served	„Zuerst gekommen – zuerst bedient“ (Bedienreihenfolge)
FMS	Flight Management System	Flugmanagementsystem
FRA		Flughafen Frankfurt am Main
GA (G/A)	Go-Around	Durchstartmanöver
HMI	Human-Machine Interface	Mensch-Maschine-Schnittstelle
ICAO	International Civil Aviation Organization	Internationale Zivilluftfahrt-Organisation
KI		Künstliche Intelligenz
KT	knot(s)	Knoten (Seemeilen pro Stunde)
LSCP	Latest Safe Clearance Point	Entspricht dem im Text verwendeten Begriff „Gate“
ML	Machine Learning	Maschinelles Lernen
MUC		Flughafen München
NM	nautical mile	Nautische Meile (1 NM = 1,852 km)
OODA	Observe-Orient-Decide-Act	Entscheidungszyklus nach Boyd
R/T	radiotelephony	Sprechfunk (Radiotelefonie)

RECAT-EU	Wake Turbulence Re-categorisation (EU)	Neu-Kategorisierung der Wirbelschleppengruppen (EU)
ROT	Runway Occupancy Time	Pistenbelegungszeit (Zeit bis „Piste frei“)
RWY	Runway	Piste
SESAR	Single European Sky ATM Research	Forschungs- und Umsetzungsinitiative für den Einheitlichen Europäischen Luftraum
SESAR JU	SESAR Joint Undertaking	SESAR-Gemeinschaftsunternehmen der EU
SID	Standard Instrument Departure	Standardabflugverfahren
STAR	Standard Terminal Arrival Route	Standard-Anflugverfahren
TMA	Terminal Manoeuvring Area	An- und Abflugkontrollbereich
WTC	Wake Turbulence Category	Wirbelschleppenkategorie

Glossar – spezialisierte Begriffe (ATM)

Begriff	Kurzdefinition
Ampel / Quantil HMI	Einfache Anzeige (grün/gelb/rot) kombiniert mit p-Werten zur schnellen Einschätzung von Exzedenzwahrscheinlichkeiten.
Calibrated Trust	Situationsadäquates Vertrauen in Assistenzsysteme mit nachweislich kalibrierten Wahrscheinlichkeiten und transparentem Gültigkeitsbereich.
Exzedenzwahrscheinlichkeit	Wahrscheinlichkeit, dass ein Grenzwert überschritten wird (z. B. „Freigabe kommt nach Gate“). Kerngröße für risikobasierte Puffer.
FMS-ladbare Freigabe	Digitale Uplink-Freigabe direkt ins FMS/AFCS (z. B. Heading/Speed/Level). Reduziert Medienbruch und beschleunigt die wirksame Umsetzung.
Frequenz Blocking	Temporäre Belegung des Sprechfunkkanals (z. B. Initial Calls, lange Readbacks), die Freigaben verzögert; modelliert als stochastischer Ressourcenengpass.

Gate Logik	Regel: Liegt bis zu einem definierten Gate (Zeit/Distanz/Höhe vor der Schwelle) keine Landefreigabe vor, initiiert die Crew den Go-Around. Synonym: Latest Safe Clearance Point (LSCP).
Headway (Anflug Headway)	Zeit-/Distanzabstand zwischen aufeinanderfolgenden Anflügen; wird operativ über Quantile und Puffer geführt.
Kognitive Latenz	Zeit für Wahrnehmen, Verstehen, Priorisieren vor der Entscheidung; variiert mit Workload und Bildaufbau.
Koinzidenz (betrieblich)	Zeitgleiches/nahezeitiges Auftreten mehrerer Einflüsse (z. B. Blocking + lange Reaktion), ohne angenommene Kausalität.
Medientrennung Voice/CPD LC	Betriebsprinzip: Zeitkritisches auf Voice, planbares Nicht-Zeitkritisches auf Data Link, mit klar definiertem Fallback.
Mixed Mode Betrieb	Abwechselnder Betrieb von Starts und Landungen auf einer Bahn; erfordert zeitlich eng gekoppelte Freigaben und robuste Gates.
Off Nominal Ereignis	Ungeplantes Ereignis (z. B. Medical, technischer Rücklauf, Sicherheitsereignis), das Varianz und Funklast kurzfristig erhöht.
Präemptives Vektorring	Vorausschauendes Eingreifen des Lotsen in die Flugführung durch Kursanweisungen (Vektoren), bevor eine Konfliktsituation tatsächlich eingetreten ist.
Quantil / p-Wert (p90/p95)	Schwellenwert, der von x % der Beobachtungen unterschritten wird. Relevanter als Mittelwerte für zeitkritische Entscheidungen (z. B. Freigabe vor Gate mit p95).
Read-back/Hear-back Fehler	Fehlinterpretation, -wiedergabe oder -wahrnehmung von Freigaben; steigt mit Botschaftslänge/Komplexität, erhöht Latenzen durch Korrekturen.
Risikobudget	Bewusst verteilte Sicherheitsmarge über die Kette Wahrnehmen/Entscheiden – Übertragen/Bestätigen – Umsetzen; wird in Quantilen geführt.
Runway Occupancy Time (ROT)	Zeitspanne bis die Piste nach Landung/Start wieder nutzbar ist; zentrale Zufallsgröße im Mixed Mode.
Stimmigkeit der Regelkreise	Harmonisierung der Loops Air Traffic Controller (ATCO) → Kommunikation → Cockpit/Flugzeug durch gemeinsame Zeitbasis, geringe Komplexität und direkte Wirksamkeit.

Supervisory Decision Maker	Rolle des Lotsen als überwachender Entscheider: prüft probabilistische Empfehlungen, priorisiert und verantwortet die Freigabe.
Tail Risiko	Seltene, aber wirkmächtige Ausreißer am Rand einer Verteilung; deterministische Planungen unterschätzen sie oft.
Trajectory Sharing (4D/i4D)	Gemeinsame Nutzung von Zeit- und Flugpfadzielen (Target Times of Arrival, 4D-Trajektorien) zwischen Boden und Bord zur Synchronisierung.

Referenzen:

[1]	Shone, R., Glazebrook, K., & Zografos, K. G. (2021). Applications of stochastic modeling in air traffic management: Methods, challenges and opportunities for solving air traffic problems under uncertainty. <i>European Journal of Operational Research</i> , 292(1), 1–26
[2]	Hoefer, C. (2003/akt. Ed.). Causal Determinism. <i>Stanford Encyclopedia of Philosophy</i> .
[3]	Eagle, A. (2022). Chance versus Randomness. <i>Stanford Encyclopedia of Philosophy</i> (Fall 2022 Ed.).
[4]	Grimmett, G. R., & Stirzaker, D. R. (2020). <i>Probability and Random Processes</i> (4. Aufl.). Oxford University Press.
[5]	Bundesgerichtshof (BGH). (2017). Urteil vom 16.05.2017 – X ZR 142/15 (Definition „von außen kommendes, auch bei äußerster Sorgfalt nicht abwendbares Ereignis“).
[6]	Diaconis, P., & Mosteller, F. (1989). Methods for Studying Coincidences. <i>Journal of the American Statistical Association</i> , 84(408), 853–861
[7]	Corver, S., & Grote, G. (2016). Uncertainty management in en-route air traffic control: A field study exploring controller strategies and requirements for automation. <i>Cognition, Technology & Work</i> , 18(3), 475–489
[8]	Poppe, M., & Buxbaum, J. (2020). Clustering Climb Profiles for Vertical Trajectory Analysis. <i>SESAR JU SID 2020</i> . – belegt starke Streuungen in Steigprofilen (operativ relevante Unsicherheit).
[9]	EUROCONTROL (2019). <i>Human Factors Integration in ATM Systems</i> (White Paper)
[10]	Bouarfa, S., Blom, H. A. P., Curran, R., & Everdij, M. H. C. (2013). Agent-based modeling and simulation of emergent behavior in air transportation. <i>Complex Adaptive Systems Modeling</i> , 1, 15.
[11]	EUROCONTROL (2021). <i>SMART Wx Task Force – WP1 Report: Collab. Best Practices</i> .
[12]	Röhner, P., Thoma, C., Schneider, W., Rohn, M., & Bott, A. (2010). Development of a low visibility forecast tool for Munich Airport.

	International Conference on Fog, Fog Collection and Dew.
[13]	SESAR Joint Undertaking & Clean Aviation. (2025, Juni). <i>Aviation Research & Innovation Strategy (ARIS)</i> . SESAR JU. Verfügbar unter: https://www.sesarju.eu/sites/default/files/documents/reports/2025_CAJU-SesarJU_ARIS_Report_FINAL.pdf
[14]	SESAR Joint Undertaking, European Commission, & EUROCONTROL. (2024). <i>European ATM Master Plan – 2025 Edition</i> . SESAR JU. Verfügbar unter: https://www.sesarju.eu/sites/default/files/documents/reports/SESAR_Master_Plan_2025.pdf
[15]	Hao, L., & Hansen, M. (2013). Scheduled block times and flight delays: Analysis of U.S. domestic flights. In <i>Proceedings of the 10th USA/Europe Air Traffic Management R&D Seminar (ATM2013)</i> , Chicago, USA. Verfügbar unter: https://atmseminar.org/seminarContent/seminar10/papers/180-Hao_0112130501-Final-Paper-4-29-13.pdf
[16]	Wang, T., Hansen, M., & Zou, B. (2019). Effects of schedule padding on delay propagation and flight cancellations. <i>Transportation Research Part A: Policy and Practice</i> , 130, 244–257. https://doi.org/10.1016/j.tra.2019.09.014
[17]	International Civil Aviation Organization (ICAO). (2016). <i>Procedures for Air Navigation Services – Air Traffic Management (PANS-ATM)</i> , Doc 4444 (16th ed., Amdt. 9). Montreal: ICAO. (Kaufpflichtig im ICAO Store). Ergänzend siehe Skybrary: https://skybrary.aero/articles/wake-turbulence-separation-minima
[18]	EUROCONTROL. (2023, März 14). <i>RECAT-EU: European Wake Turbulence Categorisation and Separation Minima on Approach and Departure (Edition 2.0)</i> . Brüssel: EUROCONTROL. Verfügbar unter: https://www.eurocontrol.int/publication/recat-eu-pws-edition-20
[19]	EUROCONTROL (2015). <i>Standard Inputs for Economic Analyses (Cost of Delay Values v4.1)</i> . Brüssel: EUROCONTROL. Verfügbar unter: https://www.eurocontrol.int/publication/standard-input-economic-analyses
[20]	EUROCONTROL (2025). <i>Performance Review Report 2024</i> . Brüssel: EUROCONTROL. Verfügbar unter: https://www.eurocontrol.int/publication/performance-review-report-2024
[21]	Wüstenbecker, N., Renkhoff, J., Zeppenfeld, D., Jameel, M., & Schier-Morgenthal, S. (2025). Analysis and prediction of pilot response times to air traffic control clearances. German Aerospace Center (DLR), Institute of Flight Guidance.
[22]	Roosens, J. (2018). Pilot response times and runway throughput. <i>European Journal of Transport and Infrastructure Research</i> , 18(4),

	45–63. Verfügbar unter: https://ejtr.tu-delft.nl/issue/2018_04
[23]	EUROCONTROL. (2016). <i>Airport Capacity Assessment Methodology (ACAM)</i> . Brüssel: EUROCONTROL. Verfügbar unter: https://ext.eurocontrol.int/ACAM
[24]	Lindner, M. (2013). <i>Entwicklung eines Algorithmus zur Anpassung der Vorhersage von Trajektorien an spezifische Flugmanöver von Luftverkehrsgesellschaften</i> (Studienarbeit). Technische Universität Dresden, Fakultät Verkehrswissenschaften. (Unveröffentlichtes Manuskript; Kopie beim Autor verfügbar).
[25]	Prats, X., Mouillet, V., Nuic, A., Cavadini, L., Lopez Leon, J., Casado, E., & Vilaplana, M. (2019). <i>Modelling flexible thrust performance for trajectory prediction applications in ATM</i> . Brüssel: EUROCONTROL. Verfügbar unter: https://www.eurocontrol.int/sites/default/files/2019-03/modelling-flexible-thrust-performance.pdf
[26]	Fricke, H., & Schultz, M. (2009). Runway operations – A comparison of time-related operational performance at major European airports. In <i>Proceedings of the 8th USA/Europe Air Traffic Management R&D Seminar (ATM2009)</i> , Napa, USA. Verfügbar unter: https://atmseminar.org/seminarContent/seminar8/papers/71-Fricke_0112090836-Final-Paper-4-29-09.pdf
[27]	Endsley, M. R. (1995). Toward a theory of situation awareness in dynamic systems. <i>Human Factors</i> , 37(1), 32–64. https://doi.org/10.1518/001872095779049543
[28]	Rasmussen, J. (1983). Skills, rules, and knowledge; signals, signs, and symbols, and other distinctions in human performance models. <i>IEEE Transactions on Systems, Man, and Cybernetics</i> , SMC-13(3), 257–266. https://doi.org/10.1109/TSMC.1983.6313160
[29]	SESAR Joint Undertaking. (2021). <i>SESAR Solutions Catalogue 2021</i> (Solution PJ.01-01, Time-Based Separation). Brüssel: SESAR JU. Verfügbar unter: https://www.sesarju.eu/sites/default/files/documents/reports/SESAR-Solutions-Catalogue-2021.pdf
[30]	SESAR Joint Undertaking. (2019). <i>PJ.16-04 Controller Working Position HMI – Deliverable on Ontology and Command Structures</i> . SESAR 2020 Project. Projektseite: https://www.sesarju.eu/projects/PJ16-04
[31]	Helmke, H., Ohneiser, O., Mühlhausen, T., Kleinert, M., & Ehr, H. (2022). Readback error detection by ASR and understanding – Results of the HAAWAll project for Isavia's en-route airspace. In <i>Proceedings of the 12th SESAR Innovation Days (SIDs 2022)</i> , Budapest, Hungary. Verfügbar unter: https://www.sesarju.eu/sites/default/files/documents/sid/2022/paper_3.pdf