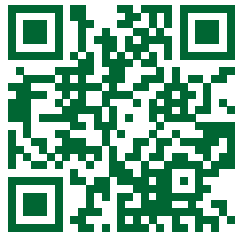


WIRTSCHAFTSPOLITIK

03 – ALLOKATIONSTHEORIE 2

Julian Hinz

Bielefeld, 23. April 2025



<https://wipo.julianhinz.com>

RÜCKBLICK: ALLOKATIONSTHEORIE 1

- Ziele staatlicher Wirtschaftspolitik (Effizienz, Gerechtigkeit)
- Bewertung von Zuständen (Pareto-Kriterium, Kaldor-Hicks)
- Soziale Wohlfahrtsfunktionen (Utilitarismus, Rawls, etc.)
- Arrow's Unmöglichkeitstheorem

HEUTE: ALLOKATIONSTHEORIE 2

- Theorie des Mechanism Design: Wie erreicht man Ziele bei dezentraler/privater Information?
- Prinzipal-Agenten-Probleme:
 - Verborgene Handlungen (Moral Hazard)
 - Private Informationen (Adverse Selection)
- Anwendung: Anreizverträge, Auktionen, öffentliche Aufträge
- Verbindung zur allgemeinen Gleichgewichtstheorie und Marktversagen

MECHANISM DESIGN

WAS IST MECHANISM DESIGN?

- Zu welchen Ergebnissen führen Regeln?
- Problem: relevante Information ist dezentral verteilt oder privat
- Regeln werden als Mechanismus bezeichnet
- Mechanismus führt zu Spiel (im Sinne der Spieltheorie), das von den Akteuren (Agenten) gespielt wird
 - Wie Spielregeln gestalten, damit Ergebnis des Spiels (im Gleichgewicht) gewünschtes Ziel entspricht?

DAS PROBLEM DER INFORMATIONSSASYMMETRIE

- vollständige Information: Planer könnte einfach anordnen, was zu tun ist
 - z.B. wer produziert, wer erhält Hilfe, wie streng sind Umweltauflagen
- Problem: Akteure haben private Informationen
 - z.B. Kosten, Präferenzen, Fähigkeiten
- Oder: Akteure führen verborgene Handlungen aus
 - nicht beobachtbar oder nicht verifizierbar/gerichtlich überprüfbar

DAS PROBLEM DER INFORMATIONSSASYMMETRIE

- Hayek (1945): Das Kernproblem einer Wirtschaftsordnung ist die sinnvolle Nutzung dezentral verteilten Wissens
 - Mechanism Design sucht Institutionen zur Informationsaggregation

DER MARKT ALS MECHANISMUS

- Marktmechanismus ist *ein* Mechanismus unter vielen möglichen
 - Regeln: Schutz von Eigentumsrechten, freier Tausch
- Mechanism Design erlaubt es, die Frage zu stellen: Gibt es bessere Mechanismen als den Markt, um bestimmte Ziele (z.B. Effizienz) unter Informationsasymmetrie zu erreichen?
- Verwandt aber anders: Allgemeine Gleichgewichtstheorie
 - untersucht Eigenschaften von Gleichgewichten *im* Marktmechanismus
 - oft unter Annahme vollständiger Information

MECHANISMEN BEI VERBORGENEN HANDLUNGEN (MORAL HAZARD)

DEFINITION: VERBORGENE HANDLUNG

- Situation in einer Prinzipal-Agenten-Beziehung
- Handlung des Agenten ist für den Prinzipal:
 - entweder nicht beobachtbar
 - oder zwar beobachtbar, aber nicht gerichtlich verifizierbar
- Konsequenz: Anreizverträge können nur auf beobachtbare und verifizierbare Ergebnisse konditionieren, die von der Handlung abhängen
- Ziel des Prinzipals: Agenten zu “wünschenswertem” Verhalten motivieren.

BEISPIEL: MITARBEITERANSTRENGUNG

- Prinzipal: Firmenbesitzer; Agent: Mitarbeiter
- Handlung: Anstrengung $e \geq 0$ des Mitarbeiters (nicht beobachtbar/verifizierbar)
- Ergebnis: Innovationserfolg (beobachtbar/verifizierbar) mit Wahrscheinlichkeit $p(e)$
- Annahmen: $p(0) = 0, p'(e) > 0, p''(e) < 0$
 - mehr Anstrengung, höhere Erfolgswahrscheinlichkeit
 - abnehmende Grenzerträge der Anstrengung

BEISPIEL: MITARBEITERANSTRENGUNG

- Gewinn bei Erfolg: π , Gewinn bei Misserfolg: 0
- Nutzen des Mitarbeiters: Lohn w minus Anstrengungskosten $e \implies U_A = w - e$
- Reservationsnutzen (Lohn in anderer Firma): $\tilde{w}$
- Anreizvertrag: Lohn bei Erfolg $\hat{w}_A$, Lohn bei Misserfolg $\check{w}_A$

AGENTENPROBLEM: OPTIMALE ANSTRENGUNG

- Gegeben Vertrag $(\hat{w}_A, \check{w}_A)$, wählt der Agent e , um erwarteten Nutzen zu maximieren:

$$\begin{aligned}\max_{e \geq 0} \quad E[U_A] &= p(e)\hat{w}_A + (1 - p(e))\check{w}_A - e \\ &= p(e)(\hat{w}_A - \check{w}_A) - e + \check{w}_A\end{aligned}$$

AGENTENPROBLEM: OPTIMALE ANSTRENGUNG

- Bedingung erster Ordnung für $e^* > 0$:

$$\frac{dE[U_A]}{de} = p'(e^*)(\hat{w}_A - \check{w}_A) - 1 = 0 \quad \Leftrightarrow \quad p'(e^*) = \frac{1}{\hat{w}_A - \check{w}_A}$$

- Interpretation: optimale Anstrengung e^* steigt mit Lohndifferenz $\Delta w = \hat{w}_A - \check{w}_A$ (da $p''(e) < 0$)
→ Höhere Anreize führen zu mehr Anstrengung

PRINZIPALPROBLEM: OPTIMALER VERTRAG

- Der Firmenbesitzer wählt den Vertrag $(\hat{w}_A, \check{w}_A)$, um seinen erwarteten Gewinn zu maximieren:

$$\max_{\hat{w}_A, \check{w}_A} E[\text{Gewinn}] = p(e^*)(\pi - \hat{w}_A) - (1 - p(e^*))\check{w}_A$$

- Nebenbedingungen:
 - Anreizkompatibilität (IC), Agent wählt e^* optimal: $p'(e^*) = 1/(\hat{w}_A - \check{w}_A)$
 - Teilnahmebedingung (PC), Agent muss mindestens Reservationsnutzen $\tilde{w}$ erhalten:

$$p(e^*)\hat{w}_A + (1 - p(e^*))\check{w}_A - e^* \geq \tilde{w}$$

→ Annahme: PC ist bindend (Prinzipal zahlt nicht mehr als nötig)

LÖSUNG DES PRINZIPALPROBLEMS (RISIKONEUTRALITÄT)

- Setze bindende PC und IC in die Zielfunktion des Prinzipals ein:

$$\begin{aligned} E[\text{Gewinn}] &= p(e^*)(\pi - \hat{w}_A) - (1 - p(e^*))\check{w}_A \\ &= p(e^*)\pi - [p(e^*)\hat{w}_A + (1 - p(e^*))\check{w}_A] \\ &= p(e^*)\pi - [\tilde{w} + e^*] \quad (\text{aus bindender PC}) \end{aligned}$$

→ äquivalent zur Maximierung der Gesamtwohlfahrt (Erwarteter Output minus Anstrengungskosten minus Reservationsnutzen)

- Ergebnis: Bei Risikoneutralität von Prinzipal und Agent führt optimaler Anreizvertrag zu Pareto-optimaler Anstrengung e^*

KOMPLIKATION: RISIKOAVERSION DES AGENTEN

- Was passiert, wenn Agent risikoavers ist (z.B. $U_A = \sqrt{w} - e$)?
- Pareto-Optimum würde erfordern, dass risikoaverser Agent vollständig gegen Einkommensrisiko versichert wird
 - erhält sicheren Lohn $\hat{w}_A = \check{w}_A$
 - Aber: sicherer Lohn bedeutet $\hat{w}_A - \check{w}_A = 0$
 - Dann ist $p'(e^*)(\hat{w}_A - \check{w}_A) - 1 = -1 < 0$, d.h. Agent wählt minimale Anstrengung $e^* = 0$

KOMPLIKATION: RISIKOAVERSION DES AGENTEN

- Trade-off: Versicherung vs. Anreize
 - Wenn Anstrengungsniveau $e^* > 0$ die Gesamtwohlfahrt maximiert, ist das Pareto-Optimum (volle Versicherung) nicht erreichbar durch einen Anreizvertrag
 - optimaler Vertrag wird einen Kompromiss eingehen: teilweise Versicherung (Lohndifferenz $\neq 0$, aber kleiner als bei Risikoneutralität) und Anstrengungsniveau, das unter First-Best-Niveau liegt
- Ineffizienz!

MECHANISMEN BEI PRIVATER INFORMATION (ADVERSE SELECTION)

DEFINITION: PRIVATE INFORMATION

- Problem des Planers: Auswahl einer Alternative x aus einer Menge X .
- Planer nicht vollständig über Präferenzen oder Eigenschaften (“Typ”) der beteiligten Individuen (Agenten) $i = 1, \dots, I$ informiert
- Private Information von Agent i : $\theta_i \in \Theta_i$
- Nutzenfunktion von Agent i : $u_i(x, \theta_i)$
 - Bewertung des Ergebnisses x hängt vom privaten Typ θ_i ab
- Ziel des Planers oft: eine effiziente oder “gute” Entscheidung x treffen, die von den (unbekannten) Typen $\theta = (\theta_1, \dots, \theta_I)$ abhängt

BEISPIELE FÜR PRIVATE INFORMATION

- Auktionen: Zahlungsbereitschaft (θ_i) der Bieter für ein Gut ist privat
- Öffentliche Aufträge: Kosten (θ_i) der anbietenden Unternehmen sind privat
- Regulierung: Kostenstruktur eines Monopolisten (θ_i) ist privat
- Versicherungen: Risikotyp (θ_i) der Versicherungsnehmer ist privat
- Sozialleistungen: Bedürftigkeit (θ_i) der Antragsteller ist privat

DEFINITION MECHANISMEN

- Ein Mechanismus ist ein Spiel, das der Planer vorgibt:
 1. Für jeden Agenten i : eine Menge von möglichen Strategien (oder Nachrichten, Aktionen) S_i
 2. Eine Ergebnisfunktion $g : S_1 \times \dots \times S_I \rightarrow X$, die jedem Profil von gewählten Strategien $s = (s_1, \dots, s_I)$ ein Ergebnis $x \in X$ zuordnet.
- Auszahlung für Agent i hängt von Typ θ_i und Ergebnis $g(s)$ ab: $u_i(g(s), \theta_i)$

BEISPIEL MECHANISMEN

- Beispiel Auktion
 - S_i = Menge möglicher Gebote
 - $g(s)$ = wer bekommt das Gut und zu welchem Preis
 - $u_i = (\text{Zahlungsbereitschaft } \theta_i - \text{Preis})$ wenn Gut erhalten, sonst 0

SOZIALE AUSWAHLFUNKTIONEN

- Eine Soziale Auswahlfunktion (“Social Choice Function”) $f : \Theta_1 \times \dots \times \Theta_I \rightarrow X$ beschreibt das *vom Planer gewünschte Ergebnis* $f(\theta)$ für jeden möglichen Vektor der privaten Informationen $\theta = (\theta_1, \dots, \theta_I)$
 - Beispiel Auktion: $f(\theta) =$ Bieter i mit höchster Zahlungsbereitschaft θ_i erhält das Gut
 - Beispiel öffentlicher Auftrag: $f(\theta) =$ Firma i mit den niedrigsten Kosten θ_i

SOZIALE AUSWAHLFUNKTIONEN

- Wären die θ_i bekannt, könnte der Planer $f(\theta)$ direkt implementieren
- Implementierungsproblem: Kann Planer Mechanismus $(S_1, \dots, S_I, g)$ finden, so dass *Gleichgewichtsergebnis* des definierten Spiels für jede mögliche Typ-Kombination θ genau $f(\theta)$ ist?
- Gleichgewichtskonzept Bayesianisches Nash-Gleichgewicht: Jeder Agent wählt Strategie $s_i^*(\theta_i)$ optimal, gegeben Erwartungen über Typen und Strategien der anderen

DIREKTE MECHANISMEN UND WAHRHEITSSAGEN

- Ein direkter Mechanismus bittet die Agenten, direkt ihren Typ θ_i zu berichten
- Der Mechanismus wendet dann eine Ergebnisregel $f(\hat{\theta})$ auf die gemeldeten Typen $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_I)$ an
- Mechanismus ist anreizkompatibel, wenn es für jeden Agenten optimal ist, Wahrheit zu sagen: $\hat{\theta}_i = \theta_i$ für alle i .
→ “Truth-telling” ist ein dann Bayessches Nash-Gleichgewicht

REVELATIONSPRINZIP

Revelationsprinzip

Eine soziale Wahlfunktion $f(\theta)$ ist genau dann durch irgendeinen Mechanismus im Gleichgewicht umsetzbar, wenn sie durch einen direkten Mechanismus **wahrheitsgemäß** umsetzbar ist.

- Enorme Vereinfachung: müssen nur direkte, anreizkompatible Mechanismen betrachten
- Gilt für bayessche Gleichgewichte, nicht notwendigerweise für dominante Strategien

BEISPIEL: WAHRHEITSGEMÄSSER MECHANISMUS IN EINER AUKTION

- Zwei Bieter mit privater Zahlungsbereitschaft $\theta_1, \theta_2 \in [0, 1]$
 - Ziel: Das Gut soll an den Bieter mit der höchsten Zahlungsbereitschaft gehen
 - Direkter Mechanismus:
 - Jeder Bieter berichtet $\hat{\theta}_i$
 - Höchstbietender bekommt das Gut
 - Preisregel: z.B. Vickrey (Zweitpreis) $\Rightarrow$ Wahrheit sagen ist optimal
- **Revelationsprinzip:** Wenn es irgendeinen Mechanismus gibt, der dieses Ziel erreicht, gibt es auch einen direkten Mechanismus, bei dem alle ehrlich sind.

MARKTVERSAGEN UND KORREKTURMÖGLICHKEITEN

HAUPTSÄTZE DER WOHLFAHRTSTHEORIE

1. Hauptsatz

Jedes Marktgleichgewicht ist Pareto-effizient.

→ unter idealisierten Bedingungen wie vollständigem Wettbewerb, keiner externer Effekte, vollständiger Information

2. Hauptsatz

Jedes Pareto-Optimum kann als Marktgleichgewicht (eventuell nach geeigneter Umverteilung der Anfangsausstattung) dargestellt werden.

→ legitimiert nicht allein die Marktwirtschaft als einzig mögliche Ordnung!

DSCHUNGELÖKONOMIE (PICCIONE & RUBINSTEIN, 2007)

- Randbemerkung: Alternative Allokationsmechanismen
- Allokation basierend auf Macht statt Märkten
- Individuen mit höherer „Macht“ verdrängen schwächere aus deren bevorzugten Objekten
- Ergebnis: stabile, Pareto-optimale Allokationen möglich, wenn Machtverteilung eindeutig
- Punkt: Auch nicht-marktliche Ordnungen können (theoretisch) effizient sein.
 - Kritik an idealisierten Annahmen (Info, Transaktionskosten) gilt oft analog

MARKTVERSAGEN BEI INFORMATIONSSASYMMETRIE

- Hauptsätze der Wohlfahrtstheorie setzen vollständige Information voraus
- oft verborgene Handlungen (Moral Hazard) und private Informationen (Adverse Selection)
 - Marktgleichgewichte nicht (First-Best) Pareto-effizient!

MARKTVERSAGEN BEI INFORMATIONSSASYMMETRIE

- Gebrauchtwagenmarkt (Akerlof's "Lemons")
 - Versicherungsmarkt (Adverse Selection, Moral Hazard)
 - Kreditmarkt (Moral Hazard, Adverse Selection)
- "Marktversagen" ... ⇒ Ruf nach staatlichen Eingriffen?

BESCHRÄNKT PARETO-OPTIMALE ERGEBNISSE

- Idee: Was ist das “bestmögliche” Ergebnis, das *überhaupt* durch irgendeinen Mechanismus (inkl. Markt) erreicht werden kann, gegeben die Informationsasymmetrie?
- Sei A die Menge aller Ergebnisse (Allokationen oder SCFs), die durch einen Mechanismus unter Berücksichtigung der Anreiz- und Teilnahmebedingungen implementierbar sind.
- Ein Ergebnis $x \in A$ heißt **beschränkt Pareto-optimal** (constrained Pareto efficient), wenn es kein anderes Ergebnis $x' \in A$ gibt, das x Pareto-dominiert.
 - man kann niemanden besser stellen, ohne anderen schlechter zu stellen, *innerhalb der Menge der unter Informationsasymmetrie erreichbaren Ergebnisse*
 - Die Menge der beschränkt Pareto-optimalen Ergebnisse $P(A)$ ist i.d.R. eine Teilmenge der (First-Best) Pareto-optimalen Ergebnisse $P(X)$.

STAATLICHER EINGRIFF: GRUNDSATZÜBERLEGUNGEN

Zentrale Frage: Wie soll der Staat auf Marktversagen reagieren?

- Verzerrung beseitigen: Staat greift direkt ein, um Ursache der Verzerrung zu eliminieren
 - z.B. Monopol zerschlagen, Informationspflichten einführen
- Verzerrung kompensieren: Staat akzeptiert Verzerrung (ggf. weil unvermeidbar), aber gleicht Folgen aus oder mildert sie
 - z.B. Pflichtversicherung, Regulierung, alternative Mechanismen

STAATLICHER EINGRIFF: GRUNDSATZÜBERLEGUNGEN

- Beispiel (Versicherungsmarkt mit Adverse Selection):
 - Problem: Nur “schlechte Risiken” versichern sich, Markt bricht zusammen
 - Option 1: Beseitigen, Infos verbessern
 - Option 2: Kompensieren, Versicherungspflicht (zwingt gute Risiken in den Pool), staatliche Versicherung, Regulierung der Verträge

THEORIE DES “SECOND BEST”

- Wenn eine der Bedingungen für das First-Best-Optimum nicht erfüllt ist (z.B. wegen Informationsasymmetrie), ist es *nicht* notwendigerweise optimal, alle anderen Bedingungen so gut wie möglich zu erfüllen
- Beispiel: Wenn Anreize wegen Moral Hazard nicht First-Best sein können, ist es vielleicht auch nicht optimal, Preise = Grenzkosten zu setzen
 - Staatliche Eingriffe müssen oft unter Kompromissbedingungen erfolgen und die Interaktion verschiedener “Verzerrungen” berücksichtigen
 - Die Analyse muss klären, ob ein staatlicher Eingriff zu einem *beschränkt* Pareto-optimalen Ergebnis führt und ob dieses dem Marktergebnis überlegen ist

ZUSAMMENFASSUNG

ZUSAMMENFASSUNG

- Mechanism Design analysiert, wie Ziele bei Informationsasymmetrie erreicht werden können.
- Prinzipal-Agenten-Theorie behandelt:
 - Verborgene Handlungen (Moral Hazard): Trade-off zwischen Anreizen und Versicherung
 - Private Information (Adverse Selection): Informationsrenten nötig, um Wahrheit herbeizuführen Revelationsprinzip vereinfacht Analyse: Fokus auf direkte, wahrheitsgemäße Mechanismen
- Anreizkompatibilität und Teilnahmebedingungen zentrale Restriktionen
- Marktversagen vs. Sinnhaftigkeit staatlicher Eingriffe