



UNIVERSITÄT
ZU KÖLN



C-SEB

CENTER FOR SOCIAL
AND ECONOMIC BEHAVIOUR
UNIVERSITY OF COLOGNE



ECONtribute

MARKETS & PUBLIC POLICY



Foto: Guentel Albers — Shutterstock.com

JAHRESTAGUNG DER GfeW IN KÖLN

GfeW

Gesellschaft für experimentelle
Wirtschaftsforschung e.V.

25. – 27. SEPTEMBER 2024

Inhaltsverzeichnis

Vorwort	1
Wichtige Orte und Anreise	2
Anreise zum Konferenzort	2
Campusplan	4
Anreise zum Konferenzdinner	5
Anreise zur Altstadtführung	6
Programmübersicht und Ablaufplan	7
Keynotes & W-LAN	8
Übersicht der Parallelsessions	9
Sessionplan Mittwoch (Session 1)	10
Sessionplan Donnerstag (Sessions 2-4)	11
Sessionplan Freitag (Sessions 5-6)	14
Teilnehmendenliste	16
Organisationsteam	20
Liste der Abstracts	21
Mittwoch (Session 1)	21
Donnerstag (Sessions 2-4)	27
Freitag (Sessions 5-6)	51

Vorwort

Liebe Teilnehmerinnen und Teilnehmer der GfeW-Jahrestagung 2024,

wir freuen uns sehr, Sie an der Universität zu Köln (UzK) herzlich willkommen zu heißen! Die Universität zu Köln zählt seit ihrer Gründung im Jahr 1388 zu den ältesten und größten Universitäten Deutschlands und verbindet auf einzigartige Weise jahrhundertealte Tradition mit einer lebendigen, innovativen Gegenwart.

Experimentelle Wirtschaftsforschung hat an der UzK eine über 20-jährige Tradition. Mit dem „Center for Social and Economic Behavior“ (C-SEB) gibt es hier ein interdisziplinäres Forschungszentrum für Forscherinnen und Forscher, die experimentell arbeiten. Auch im Exzellenzcluster ECONtribute der Universitäten Bonn und Köln ist die experimentelle Forschung prominent vertreten.

Bei der diesjährigen Tagung haben wir mit „**Künstlicher Intelligenz und Verhaltensökonomie**“ ein hochaktuelles Schwerpunktthema, zu dem wir zwei herausragende Gastredner begrüßen dürfen: Prof. Dr. Moritz Hardt vom MPI for Intelligent Systems und Prof. Dr. Iyad Rahwan vom MPI for Human Development. Insgesamt erwarten wir über 100 Teilnehmende und etwa 80 spannende Fachvorträge aus allen Bereichen der experimentellen Wirtschaftsforschung.

Wir freuen uns auf bereichernde Tage in Köln und wünschen Ihnen allen interessante Vorträge, anregende Kommentare, lebhafte Diskussionen, wertvolle Gespräche und eine rundum gelungene gute Zeit.

Herzlichst,

Bernd Irlenbusch und das Organisationsteam

Wichtige Orte und Anreise

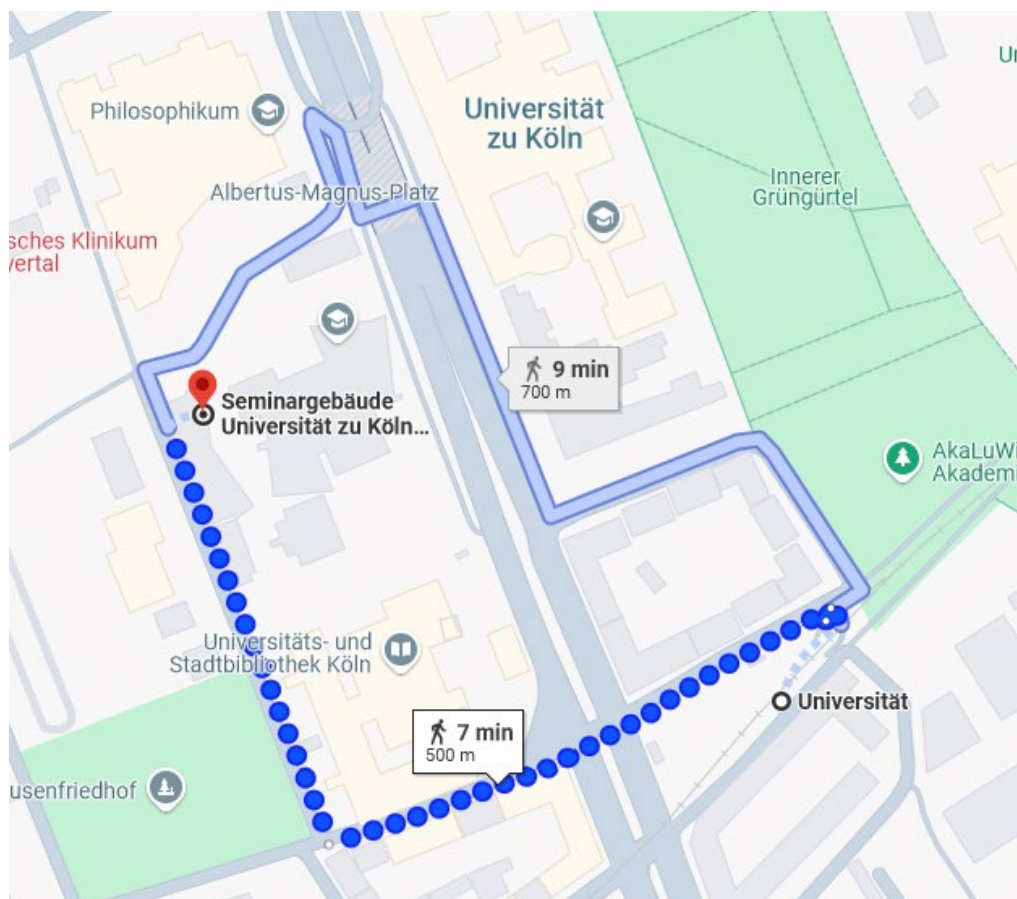
Anreise zum Konferenzort

Anreise vom Kölner Hauptbahnhof

Nehmen Sie die U-Bahn-Linie 16 oder 18 in Richtung „Bad Godesberg“ bzw. „Bonn“ bis zur Haltestelle "Neumarkt". Steigen Sie dort in die Straßenbahnlinie 9 in Richtung „Sülz“ um und fahren Sie bis zur Haltestelle „Universität“. Von dort sind es wenige Gehminuten bis zum Tagungsort.

Ticket-Automaten gibt es sowohl auf dem Bahnsteig als auch in der Bahn. Die Fahrt dauert ungefähr 20 Minuten.

Eine Taxifahrt vom Hauptbahnhof zur Universität kostet knapp 20 Euro.



Wichtige Orte und Anreise

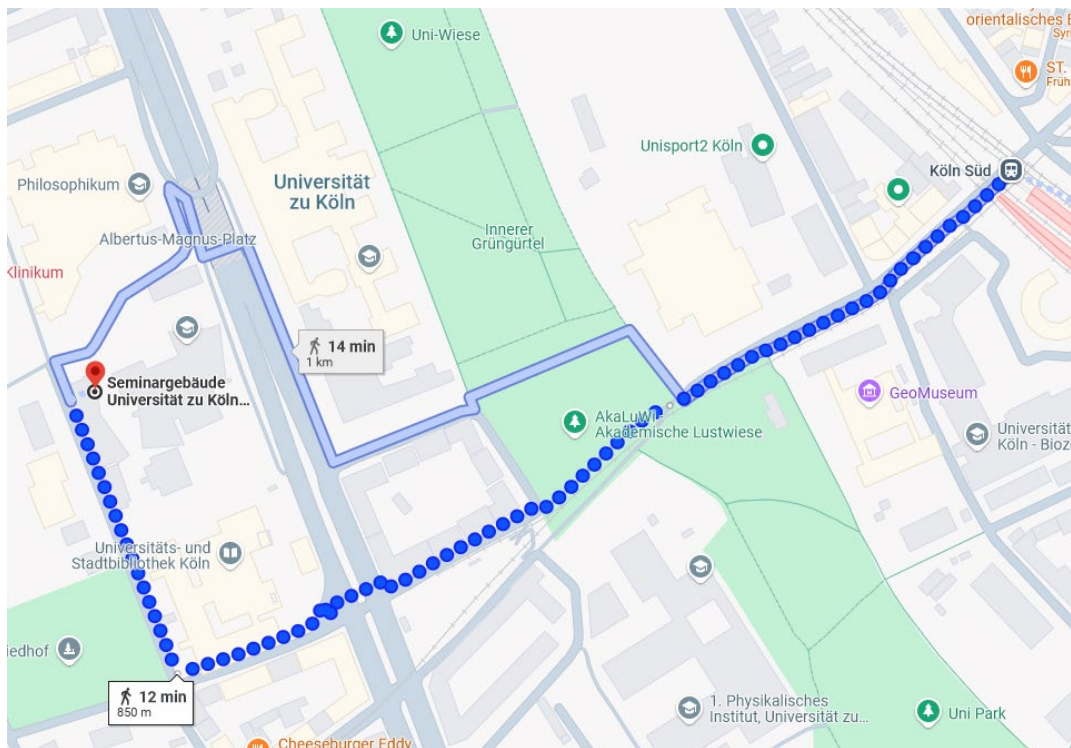
Anreise vom Bahnhof Köln Messe/Deutz

Fahren Sie mit der Straßenbahn Linie 9 in Richtung "Sülz" bis zur Haltestelle "Universität".

Eine Taxifahrt vom Bahnhof Messe/Deutz zur Universität kostet etwa 20 Euro.

Anreise vom Bahnhof Köln Süd (fußläufig)

Verlassen Sie den Bahnhof durch den Ausgang „Zülpicher Straße“. Gehen Sie nach links und folgen Sie der Zülpicher Straße in Richtung Lindenthal bis zur Kreuzung "Universitätsstraße". Biegen Sie dort rechts ab und gehen weiter bis zum Tagungsort (Gehzeit: etwa 8 – 10 Minuten).



Wichtige Orte und Anreise

Campusplan

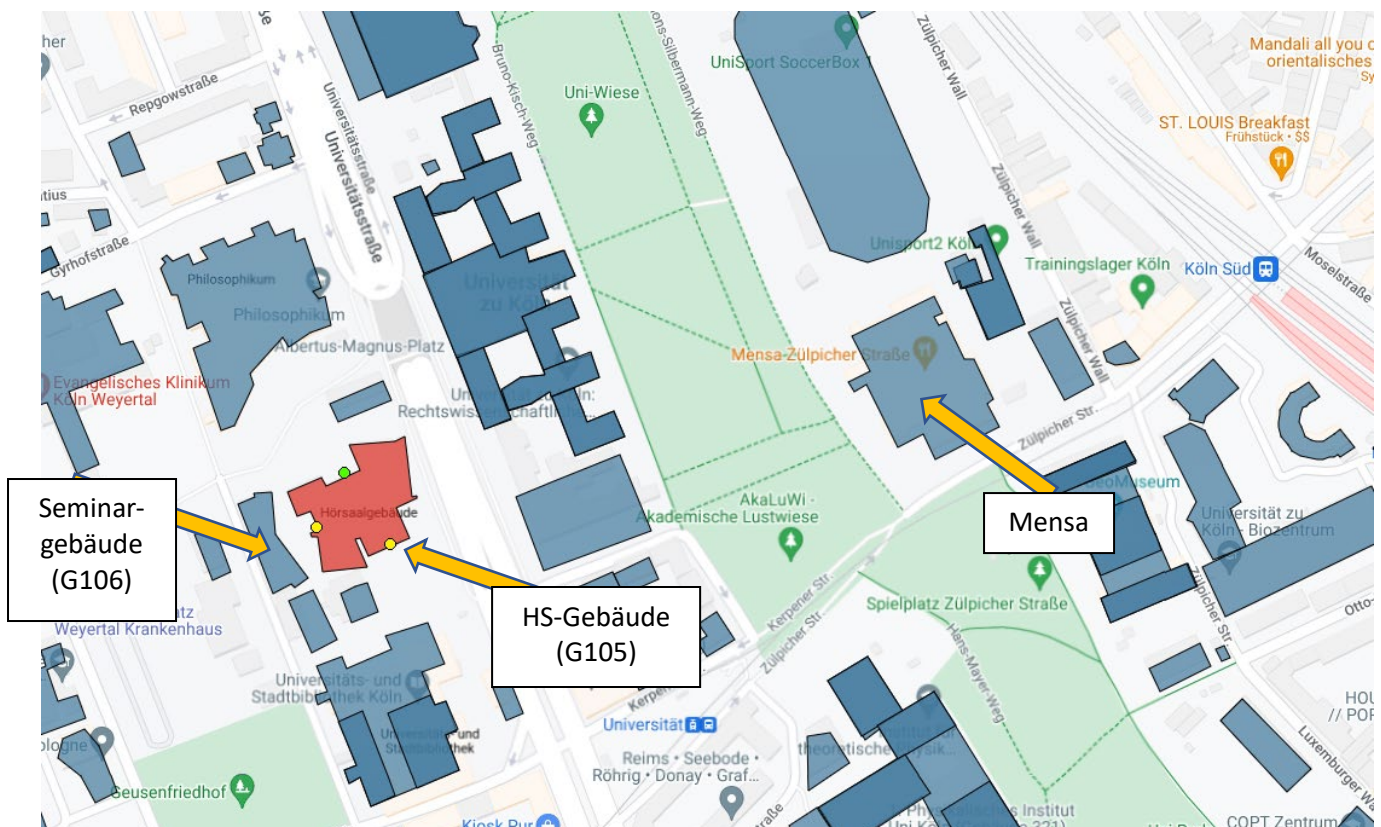
Die Jahrestagung findet auf dem Zentralcampus der Universität zu Köln statt:

Universität zu Köln

Albertus-Magnus-Platz, 50923 Köln

Hörsaalgebäude (G105) & Seminargebäude (G106)

Universitätsstraße 35 & 37, 50931 Köln



Das **Tagungsbüro** befindet sich im Erdgeschoss des Seminargebäudes (Gebäude 106). Dort wird auch die **Begrüßung** stattfinden.

Die **Keynotes** werden im Hörsaalgebäude (Gebäude 105) abgehalten. Die **Sessions** finden im 1. Stock des Seminargebäudes (Gebäude 106) statt, wobei die Namen der Räume im Seminargebäude mit dem Buchstaben „S“ beginnen.

Wichtige Orte und Anreise

Anreise zum Konferenzdinner

Das Konferenzdinner findet am Donnerstag, den 26.09.2024,
an Bord des Schiffs „**RheinHarmonie**“ statt.

Die Abfahrt erfolgt von Anlegestelle Köln 7 (Trankgassenweft), unterhalb der
Hohenzollernbrücke, in der Nähe des Kölner Hauptbahnhofs.



Copyright: KD Deutsche Rheinschiffahrt GmbH

Der Einstieg ist ab 18:30 Uhr möglich.
Die Rundfahrt und das Dinner starten
um 19:30 Uhr. Ab 22:30 Uhr liegt das
Boot wieder im Hafen und die
Veranstaltung endet um 24:00 Uhr.

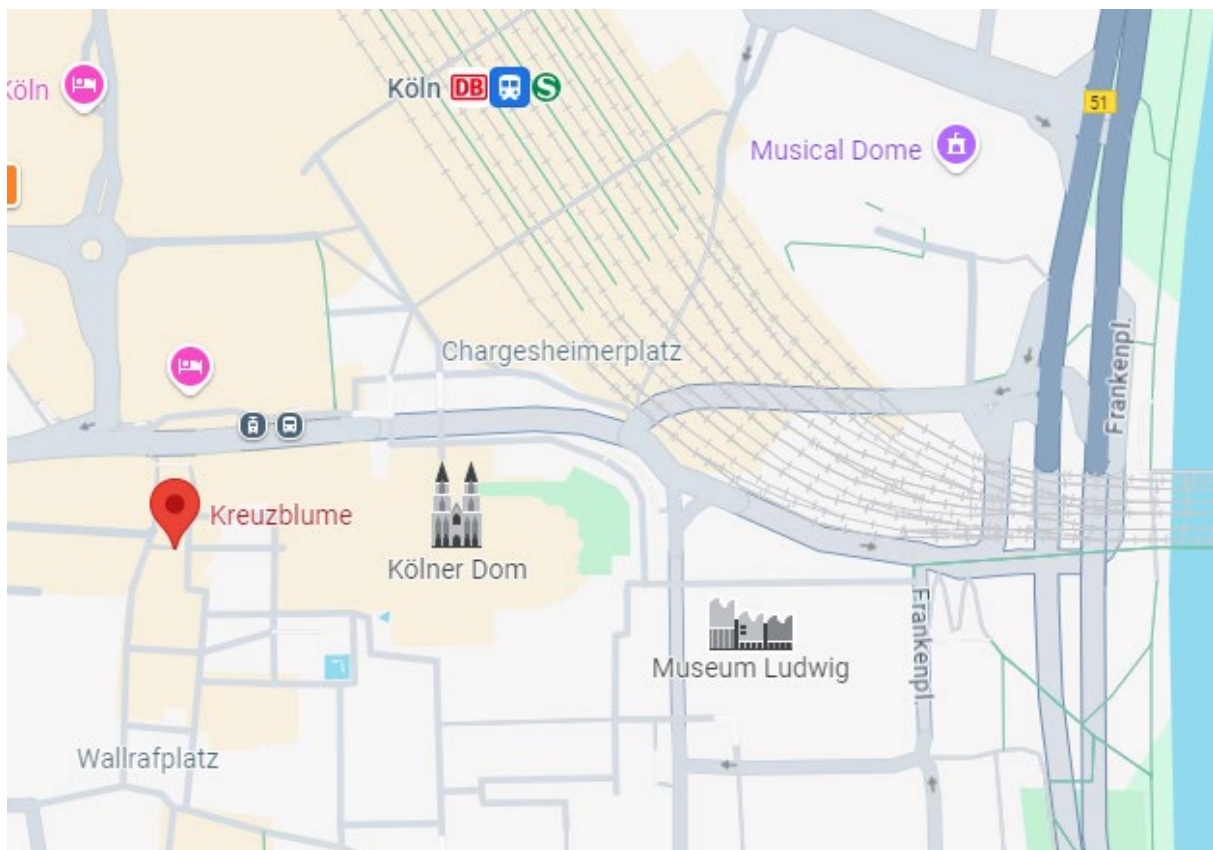


Copyright: KD Deutsche Rheinschiffahrt GmbH

Wichtige Orte und Anreise

Anreise zur Altstadtführung

Die Altstadtführung findet am 27.09.2024 um 14:30 Uhr statt. Der Treffpunkt ist an der „Kreuzblume“ (ca. 10 Meter hohes Modell der Domturmspitzen aus Beton) vor dem Kölner Dom. Die Kreuzblume steht zwischen dem Eingang von VisitKöln und dem Haupteingang des Kölner Doms, an der Domplatte auf dem Kardinal-Höffner-Platz.



Für die Anreise nutzen Sie die Linie 9 ab „Universität“ (in Richtung „Königsforst“) bis zur Haltestelle „Heumarkt“ und gehen etwa 10 Minuten zu Fuß bis zum Dom. Alternativ fahren Sie vom Heumarkt mit der Linie 5 (Richtung „Am Butzweilerhof“) bis zur Haltestelle „Dom“.

Programmübersicht

Mittwoch, 25. September

		Ort
ab 12.00 Uhr	Registrierung	Tagungsraum (Seminargebäude)
13.30 – 13.45 Uhr	Eröffnung / Grußworte	Tagungsraum (Seminargebäude)
14.00 – 15.30 Uhr	Parallelsession 1	S11-S15 (Seminargebäude)
15.30 – 16.00 Uhr	Kaffeepause	Pausenbereiche im EG (Seminargeb.)
16.00 – 17.30 Uhr	Keynote: The power of predictions Prof. Dr. Moritz Hardt, MPI for Intelligent Systems	HS A2 (Hörsaalgebäude)
17.45 – 18.45 Uhr	Mitgliederversammlung	S11 (Seminargebäude)
18:45 – 20.00 Uhr	Get-together	Tagungsraum (Seminargebäude)

Donnerstag, 26. September

09.00 – 10.30 Uhr	Parallelsession 2	S11-S16 (Seminargebäude)
10.30 – 11.00 Uhr	Kaffeepause	Pausenbereiche im EG (Seminargeb.)
11.00 – 12.30 Uhr	Parallelsession 3	S11-S16 (Seminargebäude)
12.30 – 14.00 Uhr	Mittagspause (Lunchpakete vor Ort, sonst Selbstverpflegung)	Pausenbereiche im EG (Seminargeb.)
	Forscherinnen-Netzwerktreffen	S11 (Sem.gebäude)
14.00 – 15.30 Uhr	Parallelsession 4	S11-S16 (Seminargebäude)
15.30 – 16.00 Uhr	Kaffeepause	Pausenbereiche im EG (Seminargeb.)
16.00 – 17.30 Uhr	Keynote: Opportunities at the intersection of AI and behavioral science Prof. Dr. Iyad Rahwan, MPI for Human Development	HS A2 (Hörsaalgebäude)
ab 18.30 Uhr	Konferenzdinner (Teilnahme nur nach vorheriger Anmeldung)	An Bord des Schiffs RheinHarmonie

Freitag, 27. September

09.00 – 10.30 Uhr	Parallelsession 5	S11-S15 (Seminargebäude)
10.30 – 11.00 Uhr	Kaffeepause	Pausenbereiche im EG (Seminargeb.)
11.00 – 12.30 Uhr	Parallelsession 6	S11-S14 (Seminargebäude)
12:30 – 14:00 Uhr	Verabschiedung (Lunchpakete)	Tagungsraum (Seminargebäude)
14:30 – 16:00 Uhr	Altstadtführung (Teilnahme nur nach vorheriger Anmeldung)	Treffpunkt „Kreuz- blume“ am Dom

Keynotes & W-LAN

Keynotes

Mittwoch, 25. September, 16.00 – 17.30 Uhr, HS A2

Prof. Dr. Moritz Hardt, Director, Social Foundations of Computation, Max Planck Institute for Intelligent Systems:

THE POWER OF PREDICTIONS

Donnerstag, 26. September, 16.00 – 17.30 Uhr, HS A2

Prof. Dr. Iyad Rahwan, Director, Center for Humans & Machines, Max Planck Institute for Human Development

OPPORTUNITIES AT THE INTERSECTION OF AI AND BEHAVIORAL SCIENCE

W-LAN

Für Ihre Tagungsteilnahme steht Ihnen das WLAN-Netz "Conference@University of Cologne" zur Verfügung. Die tagesaktuellen Zugangsdaten erhalten Sie von der Tagungsorganisation.

Falls Sie eduroam-Zugang über Ihre Heimatinstitution haben, können Sie alternativ auch eduroam verwenden. Informationen zur Nutzung von eduroam an der Universität zu Köln finden Sie auf folgender Webseite unter „WLAN-Zugang für Gäste mit eduroam“:

<https://rrzk.uni-koeln.de/internetzugang-web/netzzugang/wlan>

Übersicht Parallelsessions

Session	Zeit/Raum	S11	S12	S14	S15	S16
Mittwoch, 25.09.						
1	14.00 – 15:30	Risk	AI and Strategic Interaction I	Discrimination	Cooperation and Prosociality I	
Donnerstag, 26.09.						
2	09.00 – 10:30	Ethics and Fairness I	AI and Strategic Interaction II	AI and Language	Cooperation and Prosociality II (Sessionsprache: Deutsch)	Effort & Performance
3	11.00 – 12:30	Cheating and Honesty	AI and Humans I	AI and Preferences	Trust	Digitalization (Sessionsprache: Deutsch)
4	14.00 – 15:30	Volunteering and Charity (Sessionsprache: Deutsch)	AI and Humans II	Bargaining and Auctions	Social Norms and Reference Points	Health
Freitag, 27.09.						
5	09.00 – 10:30	Ethics and Fairness II (Sessionsprache: Deutsch)	AI, Transparency and Bias	Field Experiments	Cooperation and Prosociality III	
6	11:00 – 12:30	Communication	Freedom and Authority	Teamwork		

Für die einzelnen Vorträge sind ca. 20 Minuten plus etwa 10 Minuten Diskussion eingeplant.
 Vortragssprache grundsätzlich Englisch außer in einzelnen Sessions wie oben angegeben.
 Die jeweils letzten Vortragenden einer Session achten bitte auf die Einhaltung der Vortragszeiten.

Mittwoch, 25. September (Session 1)

Raum	S11	S12	S14	S15
Zeit	Risk	AI and Strategic Interaction I	Discrimination	Cooperation and Prosociality I
14:00	Monika Burckhardt: Portfolio Rebalancing: Necessary but seldom performed	Alexander Erlei: Technological Shocks and Algorithmic Decision Aids in Credence Goods Markets	Katharina Werner: The impact of gender information on hiring decisions based on self-set performance targets	Elisa Hofmann: The Krupka and Weber Social Norm Measurement: A Robustness Test Across Incentives and Methodological Variations
14:30	Kevin Piehl: Input versus Output Contingent Tax Incentives and Venture Capital - An Experimental Analysis	Andreas Orland: Playing Prisoner's Dilemma Games with Large Language Models	Stefan Grundner: Discrimination and Quotas in the Labor Market	Sven Walther: Beyond Words: Non-Verbal Cues in Virtual Collaborations
15:00	Maximilian Floto: Inflation Expectations and Economic Preferences	Hans Theo Normann: Algorithmic Cooperation	Alexandra Seidel: The Cinderella Game	Deborah Voß: Nature Rights to Promote Sustainability: A Question of Voice or Compensation for Ecosystems?

Donnerstag, 26. September (Session 2)

Raum	S11	S12	S14	S15	S16
Zeit	Ethics and Fairness I	AI and Strategic Interaction II	AI and Language	Cooperation and Prosociality II (Sessionsprache: Deutsch)	Effort & Performance
09:00	Alisa Frey: Redistribution, Moral Hazard, and Voting by Feet: An Experiment	Dmitri Bershadskyy: ChatGPT's financial discrimination between rich and poor – misaligned with human behavior and expectations.	Jannik Greif: Do You Trust Your Own Voice?	Maximilian Andres: Payoffs, Beliefs, and Cooperation in Infinitely Repeated Games	Fabian Felser: I Know What You Did Last Summer: How Past Performance Affects Future Performance Evaluation
09:30	Christian Koch: Populist propaganda and anti-migration stances: An experimental investigation	Rosemarie Nagel: A Behavioral Taxonomy of 2x2 Games		Maximilian Kuntze: Why do People Follow an Example?	Louis Strang: The Effect of Task (Mis)Matching and Self-Selection on Intrinsic Motivation and Performance
10:00	Julian Conrads: Corporate Responsibility Increases Consumers' Willingness to Pay: Evidence from two Framed Field Experiments	Heinrich Nax: Karma: An experimental investigation of a 'closed' intertemporal token system	Florian E. Sachs: Seeing is believing: How (X)AI-augmented advice guides human task-solving behavior	Nicole Middendorf: Motivational Effects of Monetary Incentives and Social Recognition on Parental Engagement	

Donnerstag, 26. September (Session 3)

Raum	S11	S12	S14	S15	S16
Zeit	Cheating and Honesty	AI and Humans I	AI and Preferences	Trust	Digitalization (Sessionsprache: Deutsch)
11:00	Lilith Burgstaller: Prevalence and perceptions of undeclared work and tax evasion	Nina Rulié: Delegation to Pricing Algorithms: Experimental Evidence	Christian Feige: Learning to be kind: neural networks as reciprocal altruists in allocation games	Dinithi Jayasekara: Do people trust Artificial Intelligence? A Repeated Trust Game Experiment	
11:30	Nils Köbis: Ethical Risks of Delegation to AI	Julia Werner: Searching for Stars or Avoiding Catastrophes: The Role of Loss Aversion for Trust in AI	Paul M. Gorny: Do Personalized AI Predictions Change Subsequent Decision-Outcomes? Artificial versus Swarm Intelligence	Matthias Kasper: Improving Trust as a Tool for Improving Tax Compliance	Tom Reuscher: Bridging the Communication Gap: Real-Time Visualization of Eye Contact to Support Cooperation in Virtual Teams
12:00	Rainer Michael Rilke: Beliefs and Group Dishonesty: The Role of Strategic Interaction and Responsibility	Luisa Santiago Wolf: AI or human? Applicants' decisions in discrimination settings	Irenaeus Wolff: Don't do what I might choose, but please do what will be better: an algorithm's role, feedback, & algorithm reliance	Lucas Braun: Handshakes and Trust in Virtual Reality	Lara Berger: How digital media markets amplify news sentiment

Donnerstag, 26. September (Session 4)

Raum	S11	S12	S14	S15	S16
Zeit	Volunteering and Charity (Sessionsprache: Deutsch)	AI and Humans II	Bargaining	Social Norms and Reference Points	Health
14:00	Maximilian Hiller: Recruiting Episodic Volunteers: A Field Experiment on Gift-Giving vs. Requesting a Favor	Eduard Buzila: Can the Use of Large Language Models in Higher Education Empower Students to Increase their Learning Success?	Holger Rau: The Economic Effects of Remote-Bargaining vs. Face-to-Face Bargaining	Amanda März: Sanctions and Social Norms	Martin Strobel: Human Infection Risk Assessment in Social Networks: The Role of Network Characteristics
14:30	Bernd Josef Leisen: Boosting Performance Quality in Episodic Micro- Volunteering: A Collaborative Field Experiment	Yefim Roth (Ori Plonsky): Appreciation to biased algorithms, and aversion from unbiased algorithms.	Behnud Mir Djawadi: Between egalitarianism and utilitarianism: experimental evidence in unstructured bargaining	Blanca Tena Estrada: Social Networks and Coordination: A Simulation Experiment	Anna Katharina Stirner: An experiment on decision support in medicine
15:00		Yefim Roth: Overreliance, Backfire, and the Importance of (Mis)trust in Human- Automation Interactions	Dmitri Bershadskyy (Nina Ostermaier): Lie Against AI: Revealing Private Information through AI in an Economic Experiment	Till Vater: Determinants and Dynamics of Collaborative Exploitation	Stefan Traub: Sow the Doubt, Reap the Inefficiency: A Laboratory Study on Communicating Risks in the Presence of Non-Events

Freitag, 27. September (Session 5)

Raum	S11	S12	S14	S15
Zeit	Ethics and Fairness II (Sessionsprache: Deutsch)	AI, Transparency and Bias	Cooperation and Prosociality III	Field Experiments
09:00	Janina Kraus: Markets, Social Responsibility and the Replacement Logic	Anika Bittner: Literature review on Explainability of Artificial Intelligence from a Behavioral Perspective (Working title)	Fabian Hoffmann: Image Concerns and Reputation	Christian König genannt Kersting: Experimenting with Financial Professionals
09:30	Dirk Kieseewetter: The Impact of Tax Culture on Tax Rate Structure Preferences: Results from a Vignette Study with Migrants and Nonmigrants in Germany	Marius Protte: How does being involved in AI training affect expert user adherence - An experiment in the context of AutoML	Christine Meemann: Institutional Change in the Infinitely Repeated Prisoners' Dilemma	Anastasia Danilov: Do AI skills improve employment chances? Evidence from a field experiment
10:00	Heike Hennig-Schmidt: Fairness Deliberations and Fair Allocations in Symmetric and Asymmetric Bargaining An Experimental Study on Group Decisions in Germany and China	Carina Hausladen: Measuring social bias in face perception of a vision-language model	Philipp Monschau: Inflated Rules in Threshold Public Goods Games	Devin Kwasniok: Motivating the Weak: The Impact of Group-Based Incentives on Teams

Freitag, 27. September (Session 6)

Raum	S11	S12	S14
Zeit	Communication	Freedom and Authority	Teamwork
11:00	Sabrina Plaß: Experimentally examining multiple normative expectations in whistleblowing and the influence of social information interventions	Mathilde Draeger: Giving a Voice - Increasing Individual Self-Expression to Enhance Group Welfare and the Resilience to System Disbelief	Kevin Grubiak: Social Closeness in Coordination Games with Conflict
11:30	Mey-Ling Sommer: Misrepresenting Risks: A sender-receiver model for efficient outcomes	Max R. P. Grossmann: Knowledge and Freedom	Lisa Klinkenberg: Creative Performance in Teams: The Effect of Workplace Settings and Self-Selection - an Experimental Study
12:00	Marina Schröder: Avoid, recycle, compensate - An experimental study on nudging environmental behavior	Gönül Doğan: Authoritarian Preferences	Nathalie Römer: It's a match! Team formation and performance in innovation-related tasks

Teilnehmende

Name	Universität	Kontakt
Andres, Maximilian	Universität Potsdam	maximilian.andres@uni-potsdam.de
Berger, Lara	Universität zu Köln	lara.berger@uni-koeln.de
Bershadskyy, Dmitri	Otto-von-Guericke Universität, Magdeburg	dmitri.bershadskyy@ovgu.de
Betz, Dirk	TIB - Leibniz Information Centre for Science and Technology	dirk.betz@tib.eu
Biniossek, Claudia	TIB	claudia.biniossek@tib.eu
Bittner, Anika	Technische Universität Hamburg	anika.bittner@tuhh.de
Braun, Lucas	RWTH Aachen	Lucas.Braun@rwth-aachen.de
Burckhardt, Monika	Universität Mannheim	monika.burckhardt@uni-mannheim.de
Burgstaller, Lilith	Walter Eucken Institut / Universität Freiburg	burgstaller@eucken.de
Buzila, Eduard	Otto-von-Guericke Universität Magdeburg	Eduard.Buzila@ovgu.de
Conrads, Julian	Technische Hochschule Mittelhessen	julian.conrads@w.thm.de
Danilov, Anastasia	Humboldt-Universität zu Berlin	anastasia.danilov@hu-berlin.de
Das, Tanmoy	Georg-August Universität Göttingen	tanmoy.das@uni-goettingen.de
Doğan, Gönül	Universität zu Köln	dogan@wiso.uni-koeln.de
Draeger, Mathilde	Otto-von-Guericke- Universität Magdeburg	mathilde.draeger@ovgu.de
Elrich, Alina	Universität Paderborn	alina.elrich@uni-paderborn.de
Erlei, Alexander	Georg-August Universität Göttingen	alexander.erlei@wiwi.uni-goettingen.de
Feige, Christian	Unabhängiger Forscher	Christian-feige@web.de
Felser, Fabian	TU Dresden	fabian.felser@tu-dresden.de
Floto, Maximilian	Leibniz Universität Hannover	floto@gif.uni-hannover.de
Frank, Björn	Universität Kassel	frank@uni-kassel.de
Frey, Alisa	Heinrich-Heine- Universität Düsseldorf	alisa.frey@hhu.de
Geschwind, Stephan	Universität Passau	stephan.geschwind@uni-passau.de

Teilnehmende

Name	Universität	Kontakt
Gorny, Paul M.	Karlsruher Institut für Technologie	paul.gorny@kit.edu
Graf Lambsdorff, Johann	Universität Passau	jlambsd@uni-passau.de
Greif, Jannik	Otto-von-Guericke-Universität Magdeburg	jannik.greif@ovgu.de
Grossmann, Max R. P.	Universität zu Köln & Mercatus Center, George Mason University	m.grossmann@uni-koeln.de
Grubiak, Kevin	Universität Passau	kevin.grubiak@uni-passau.de
Grundner, Stefan	Universität zu Köln	stefan.grundner@uni-koeln.de
Gürerk, Özgür	Universität zu Köln	guererk@wiso.uni-koeln.de
Harbring, Christine	RWTH Aachen	christine.harbring@org.rwth-aachen.de
Hausladen, Carina	ETH Zurich	carina.hausladen@gess.ethz.ch
Heinrich, Kai	Otto-von-Guericke-Universität, Magdeburg	kai.heinrich@ovgu.de
Hennig-Schmidt, Heike	Universität Bonn	hschmidt@uni-bonn.de
Hiller, Maximilian	Universität Vechta	maximilian.hiller@uni-vechta.de
Hoffmann, Fabian	Universität zu Köln	f.hoffmann@wiso.uni-koeln.de
Hofmann, Elisa	Friedrich-Schiller-Universität Jena	elisa.hofmann@uni-jena.de
Hörrmann, Michelle	Karlsruher Institut für Technologie	michelle.hoerrmann@kit.edu
Ilgün, Can	RWTH Aachen	can.ilguen@rwth-aachen.de
Irlenbusch, Bernd	Universität zu Köln	bernd.irlenbusch@uni-koeln.de
Jayasekara, Dinithi	Singapore University of Technology and Design	dinithi@sutd.edu.sg
Kasper, Matthias	Walter Eucken Institut	kasper@eucken.de
Kerkhoff, Martin	Otto-von-Guericke-Universität, Magdeburg	martin.kerkhoff@ovgu.de
Keser, Claudia	Georg-August-Universität Göttingen	ckeser@uni-goettingen.de
Kiesewetter, Dirk	Julius-Maximilians-Universität Würzburg	dirk.kiesewetter@uni-wuerzburg.de
Klinkenberg, Lisa	RWTH Aachen University	lisa.klinkenberg@org.rwth-aachen.de
Köbis, Nils	Universität Duisburg-Essen	nils.kobis@uni-due.de
Koch, Christian	Universität Wien	chris.koch@univie.ac.at

Teilnehmende

Name	Universität	Kontakt
Koch, Marcel	RWTH Aachen University	marcel.koch@org.rwth-aachen.de
König genannt Kersting, Christian	Universität Innsbruck	christian.koenig@uibk.ac.at
Kraus, Janina	TU Clausthal	janina.kraus@tu-clausthal.de
Kuntze, Maximilian	Universität Vechta	maximilian.kuntze@uni-vechta.de
Kwasniok, Devin	Universität Vechta	Devin.Kwasniok@uni-vechta.de
Lauer, Thomas	Universität Erfurt	thomas.lauer@uni-erfurt.de
Leisen, Bernd Josef	Universität Vechta	bernd.josef.leisen@uni-vechta.de
Lutsenko, Tetiana	Universität Paderborn	lutsenko@mail.uni-paderborn.de
März, Amanda	Walter Eucken Institut	maerz@eucken.de
Meemann, Christine	Helmut-Schmidt-Universität Hamburg	meemannc@hsu-hh.de
Middendorf, Nicole	Universität Vechta	nicole.middendorf@uni-vechta.de
Mir Djawadi, Behnud	Universität Paderborn	behnud.mir.djawadi@uni-paderborn.de
Monschau, Philipp	RWTH Aachen	philipp.monschau@hrm.rwth-aachen.de
Nagel, Rosemarie	ICREA-UPF-BSE	rosemarie.nagel@upf.edu
Nax, Heinrich	ETH/UZH	hnax@ethz.ch
Nieken, Petra	Karlsruher Institut für Technologie (KIT)	petra.nieken@kit.edu
Normann, Hans Theo	Heinrich-Heine-Universität Düsseldorf	normann@hhu.de
Orland, Andreas	Corvinus University of Budapest	andreas.orland@uni-corvinus.hu
Piehl, Kevin	Leibniz Universität Hannover	piehl@wipol.uni-hannover.de
Plaß, Sabrina	Universität Paderborn	sabrina.plass@upb.de
Protte, Marius	Universität Paderborn	marius.protte@uni-paderborn.de
Pütz, Lisa	RWTH Aachen	lisa.puetz@hrm.rwth-aachen.de

Teilnehmende

Name	Universität	Kontakt
Rau, Holger	Universität Duisburg-Essen & Universität Göttingen	holger.rau@uni-goettingen.de
Reuscher, Tom	Karlsruher Institut für Technologie (KIT)	tom.reuscher@kit.edu
Rilke, Rainer Michael	WHU - Otto Beisheim School of Management	rainer.rilke@whu.edu
Römer, Nathalie	Leibniz Universität Hannover	roemer@wipol.uni-hannover.de
Roth, Yefim	University of Haifa	yroth1@univ.haifa.ac.il
Rulié, Nina	Heinrich-Heine-Universität Düsseldorf	rulie@dice.hhu.de
Sachs, Florian	Universität zu Köln	florian.sachs@uni-koeln.de
Santiago Wolf, Luisa	Universität zu Köln	santiago-wolf@wiso.uni-koeln.de
Schneider, Iris	TU Dresden	iris.schneider@tu-dresden.de
Schröder, Marina	Leibniz Universität Hannover	schroeder@wipol.uni-hannover.de
Seebauer, Michael	MPI zur Erforschung von Gemeinschaftsgütern	seebauer@coll.mpg.de
Seidel, Alexandra	Otto-von-Guericke Universität Magdeburg	alexandra.seidel@ovgu.de
Simons, Stella	RWTH Aachen University	stella.simons@org.rwth-aachen.de
Sommer, Mey-Ling	Helmut Schmidt Universität	sommerme@hsu-hh.de
Stirner, Anna Katharina	Universität zu Köln	stirner@wiso.uni-koeln.de
Strang, Louis	Universität zu Köln	strang@wiso.uni-koeln.de
Strobel, Martin	Maastricht University	m.strobel@maastrichtuniversity.nl
Tena Estrada, Blanca	Universität Kassel	blanca.tena@uni-kassel.de
Traub, Stefan	Helmut-Schmidt-Universität Hamburg	traubs@hsu-hh.de
Vater, Till Moritz	Maastricht University	till.vater@maastrichtuniversity.nl
Voß, Deborah	Universität Passau	deborah.voss@uni-passau.de
Walther, Sven	Karlsruher Institut für Technologie	sven.walther@kit.edu
Weimann, Joachim	Otto-von-Guericke-Universität, Magdeburg	joachim.weimann@ovgu.de
Werner, Julia	Universität Konstanz	julia.werner@uni-konstanz.de
Werner, Katharina	Universität zu Köln	katharina.werner@uni-koeln.de
Wolff, Irenaeus	Universität Konstanz	irenaeus.wolff@uni-konstanz.de

Organisationsteam

Stefan Grundner

Özgür Gürerk

Fabian Hoffmann

Bernd Irlenbusch

Seminar für Unternehmensentwicklung und Wirtschaftsethik

Jennifer Mayer

Claudia Töpper-Ko

C-SEB Center for Social
and Economic Behavior

Department Corporate Development

Wissenschaftliches Programmkomitee

Arno Apffelstaedt, Özgür Gürerk, Bernd Irlenbusch und Louis Strang

Bei Fragen wenden Sie sich bitte an:

gfew-2024@uni-koeln.de

Liste der Abstracts

Session 1 (Mittwoch 25.09., 14:00 Uhr)

S11	S12	S14	S15
Risk	AI and Strategic Interaction I	Discrimination	Cooperation and Prosociality I
Portfolio Rebalancing: Necessary but seldom performed Monika Burckhardt (Universität Mannheim)	Technological Shocks and Algorithmic Decision Aids in Credence Goods Markets Alexander Erlei (Georg-August Universität Göttingen)	The impact of gender information on hiring decisions based on self-set performance targets Katharina Werner (Universität zu Köln)	The Krupka and Weber Social Norm Measurement: A Robustness Test Across Incentives and Methodological Variations Elisa Hofmann (Friedrich-Schiller-Universität Jena)
Input versus Output Contingent Tax Incentives and Venture Capital - An Experimental Analysis Kevin Piehl (Leibniz Universität Hannover)	Playing Prisoner's Dilemma Games with Large Language Models Andreas Orland: (Corvinus University of Budapest)	Discrimination and Quotas in the Labor Market Stefan Grundner: (Universität zu Köln)	Beyond Words: Non-Verbal Cues in Virtual Collaborations Sven Walther (Karlsruher Institut für Technologie)
Inflation Expectations and Economic Preferences Maximilan Floto (Leibniz Universität Hannover)	Algorithmic Cooperation Hans Theo Normann (Heinrich-Heine-Universität Düsseldorf)	The Cinderella Game Alexandra Seidel (Otto-von-Guericke Universität)	Nature Rights to Promote Sustainability: A Question of Voice or Compensation for Ecosystems? Deborah Voß (Universität Passau)

Risk

TITEL: Portfolio Rebalancing: Necessary but seldom performed

AUTOR*INNEN: Monika Burckhardt, Martin Weber

ABSTRACT: We experimentally research portfolio rebalancing in a hand-collected sample of retail investors and financial advisers. Participants construct investment portfolios, observe portfolio weight changes, and are asked to manage the portfolios. The treatment condition is the time frame of compounded returns - three or six years. With extreme portfolio weight deviations of at least 15 percentage points (e.g., one portfolio weight increased from 15% to 30%), 60% of the retail investors and 46% of the advisers do not rebalance in our experiment. Women rebalance less than men but have a greater risk tolerance. Significant deviations from the initial asset allocation and portfolio risk arise if individuals do not rebalance, which may be undesirable and unsustainable to the investors. Our results confirm the necessity for better awareness and implementation of rebalancing as a risk-control mechanism.

TITEL: Input versus Output Contingent Tax Incentives and Venture Capital - An Experimental Analysis

AUTOR*INNEN: Kay Blaufus, Kevin Piehl, Marina Schröder

ABSTRACT: Access to venture capital is a key determinant towards fostering entrepreneurship. Thus, it is a regularly stated policy goal to increase the provision of venture capital. We suggest an experimental setup to analyze differences in the effectiveness of input contingent incentives (i.e., tax reductions contingent on venture capital investments) and output contingent incentives (i.e., tax reductions for profits generated through venture capital investments). We propose a sequential experimental approach where venture capitalists first decide on investing in a venture. Contingent on investing, venture capitalists then decide on providing real-effort to improve the venture. While venture capitalists in our design can improve the venture, they cannot change all aspects of the venture. We plan to analyze differences in the effect of input and output contingent incentives on the behavior of venture capitalists both in the investment stage and in the real-effort stage. The generated insights are going to contribute to a better understanding of interventions to foster the provision of venture capital.

TITEL: Inflation Expectations and Economic Preferences

AUTOR*INNEN: Lena Dräger, Maximilian Floto, Marina Schröder

ABSTRACT: This paper investigates the influence of time and risk preferences on the processing of uncertain inflation-related information and the formation of inflation expectations among households and experts. A randomized controlled trial (RCT) approach was employed to conduct an online survey of 3,266 German households and 663 employees from German banks and insurance companies. The findings indicate that risk preferences significantly influence the processing of information with macroeconomic

uncertainty and the formation of inflation expectations. On average, risk-averse households and experts have higher inflation expectations than those who are risk-tolerant. However, while we do not find differences in inflation expectation adjustment between the two risk preferences, we do find that the uncertainty of risk-tolerant households and experts is reduced more by inflation forecasts. Furthermore, we observe that there is a distinction between information with and without macroeconomic uncertainty. The impact of time preferences on the processing of information and the formation of inflation expectations remains less clear. On average, impatient households have higher inflation expectations and therefore revise their prior beliefs more strongly. However, this is not reflected in the formation of inflation expectations.

AI and Strategic Interaction I

TITEL: Technological Shocks and Algorithmic Decision Aids in Credence Goods Markets

AUTOR*INNEN: Alexander Erlei, Lukas Maub

ABSTRACT: In credence goods markets such as health care or repair services, consumers rely on experts with superior information to adequately diagnose and treat them. Experts, however, are constrained in their diagnostic abilities, which hurts market efficiency and consumer welfare. Technological breakthroughs that substitute or complement expert judgments have the potential to alleviate consumer mistreatment. This article studies how competitive experts adopt novel diagnostic technologies when skills are heterogeneously distributed and obfuscated to consumers. We differentiate between novel technologies that increase expert abilities, and algorithmic decision aids that complement expert judgments, but do not affect an expert's personal diagnostic precision. We show that high-ability experts may be incentivized to forego the decision aid in order to escape a pooling equilibrium by differentiating themselves from low-ability experts. Results from an online experiment support our hypothesis, showing that high-ability experts are significantly less likely than low-ability experts to invest into an algorithmic decision aid. Furthermore, we document pervasive under-investments, and no effect on expert honesty.

TITEL: Playing Prisoner's Dilemma Games with Large Language Models

AUTOR*INNEN: Andreas Orland, Kazuhiro Takemoto

ABSTRACT: We present the findings from an experiment in which three versions of ChatGPT-3.5-turbo make decisions in a one-shot Prisoner's Dilemma. Our study varies four key dimensions: (i) the game's payoff parameters, (ii) the number of simultaneous interaction partners for each player, (iii) the strategy space available to the players, and (iv) the sequence in which decisions and beliefs are elicited. Our results show that variations in payoff parameters do not influence the LLM's simulated behavior. However, changes in the other dimensions affect cooperation rates in a manner consistent with previous human-based experiments.

TITEL: Algorithmic Cooperation

AUTOR*INNEN: Bernhard Kasberger, Simon Martin, **Hans-Theo Normann**, Tobias Werner

ABSTRACT: Algorithms play an increasingly important role in economic situations. Often these situations are strategic, where the artificial intelligence may or may not be cooperative. We study the determinants and forms of algorithmic cooperation in the infinitely repeated prisoner's dilemma. We run a sequence of computational experiments, accompanied by additional repeated prisoner's dilemma games played by humans in the lab. We find that the same factors that increase human cooperation largely also determine the cooperation rates of algorithms. However, algorithms tend to play different strategies than humans. Algorithms cooperate less than humans when cooperation is very risky or not incentive compatible.

Discrimination

TITEL: The impact of gender information on hiring decisions based on self-set performance targets

AUTOR*INNEN: Susanna Grundmann, Bettina Rockenbach, **Katharina Werner**

ABSTRACT: Labor market inequalities such as the gender wage gap and the underrepresentation of women in leadership positions persist until today. Potential reasons for these inequalities include employers' discrimination against women in the hiring process, but also gender differences on the employee side, as, for example, in risk seeking, altruism and competitiveness. A potential additional cause that has received less attention up to now are gender differences in self-promotion and target-setting. In many settings, self-set targets are explicitly or implicitly used as performance indicators and thus as the basis for hiring or promotions to higher-level positions. In this paper, we study an experimental labor market in which employee candidates announce a self-set performance target for their work on a real effort task. In light of these targets, employers then hire an employee and the hired employee's actual performance determines the employer's and the employee's payoff. Our 2x2 design varies (1) whether or not employers are informed about the employee candidates' genders, and (2) the severity of the employer's payoff consequences of the hired employee missing their performance target. This design allows us to draw causal inference on the interaction effects of gender-anonymity in the hiring process and the payoff relevance of the performance target. We find that women, given equal ability, set lower targets than men. Higher targets increase the chance of being hired, while employers' expected target-performance gap has a negative influence on hiring, in particular when missing the target has severe payoff consequences for the employer. Gender-revealing applications increase women's chances of being hired, resulting in significantly higher payoffs for women. With gender information, employers seem to correct the self-set targets, as they expect a larger target-performance gap for men than for women. We thus conclude that gender-anonymous applications may backfire, as they may preclude employers from taking (women-favoring) gender-stereotypes into account.

TITEL: Discrimination and Quotas in the Labor Market

AUTOR*INNEN: Oliver Gürtler, Bernd Irlenbusch, **Stefan Grundner**

ABSTRACT: We investigate discrimination in the labor market. Three employers compete for six employees after promotion signals have been revealed to the market. In a baseline treatment all employers and

employees belong to different groups. No discrimination should occur in this treatment. In a second treatment we induce same group pairs. In this treatment endogenous statistical discrimination should occur and different equilibria should be selected. In further treatments we introduce different quota rules for promotion decisions and investigate their effect on wages.

TITEL: The Cinderella Game

AUTOR*INNEN: Jannik Greif, Franziska Rumpel, Abdolkarim Sadrieh, **Alexandra Seidel**

ABSTRACT: The exclusion of an individual from group benefits can be damaging in numerous ways, including physiological and psychological dimensions. The central objective of the study is to understand the selection criteria used for exclusion decisions. We also investigate whether the information on being excluded by others amplifies or abates the likelihood of being selected for exclusion. In our online study, we invited a set of subjects about whom we collected information in multiple dimensions, including behavioral characteristics and appearance. Using the distributions of the attributes, we created fictional person cards and asked the subjects to rank them according to the expected likelihood that they will be excluded by their group. We incentivized the true ranking elicitation by paying subjects payoffs that increased with the similarity to the group's overall ranking. In the treatment, we provided information about some cards being among the top 5 of an earlier experiment's overall ranking in order to examine whether subjects follow this anchor. The study is still ongoing. Results will be shared at the conference.

Cooperation and Prosociality I

TITEL: The Krupka and Weber Social Norm Measurement: A Robustness Test Across Incentives and Methodological Variations

AUTOR*INNEN: **Elisa Hofmann**, Deliah Wagnel

ABSTRACT: The social norm measurement by Krupka and Weber (2013) employs an incentivized coordination game to elicit social norms as shared perceptions of social appropriateness. This method has gained substantial traction in behavioral economics ever since its conceptualization in 2013. Given the ongoing debate in economics and psychology on the necessity of incentives for reliable results, we conducted a robustness test of the Krupka and Weber social norm measurement focusing on incentivization and, in addition, on order of choices, and length of scale. In an online-experiment with N = 422 subjects, we find that the Krupka and Weber social norm measurement yields consistent results regardless of incentivization. Furthermore, we demonstrate that the Krupka and Weber social norm measurement is partially robust towards order of choices and is sensitive to the length of scale used. These findings suggest that while the Krupka and Weber method is robust in many respects, certain methodological factors can influence its outcomes.

TITEL: Beyond Words: Non-Verbal Cues in Virtual Collaborations

AUTOR*INNEN: Michelle Hörrmann, Petra Nieken, Sven Walther

ABSTRACT: Virtual communication has become commonplace, especially for geographically dispersed groups that collaborate on joint projects despite lacking prior relationships. Unfortunately, these groups often encounter coordination failures, resulting in low overall cooperation rates. In our study, we address this challenge by implementing a weakest-link game in a controlled large-scale experiment. Our goal is to investigate how different communication channels in virtual group meetings can mitigate these issues. Before the weakest-link game, group members participate in a virtual group meeting to get acquainted. This interaction sets the stage for subsequent collaboration. We systematically vary the availability of non-verbal cues and signals within different communication channels. These cues and signals are expected to play a crucial role in overcoming strategic uncertainty and fostering group cohesion. Depending on the treatment, the virtual group meeting is a pure audio meeting, an audio meeting with a picture of the other group members, or a video meeting. We expect that richer communication channels will lead to better coordination and a higher level of cooperation among group members. Our findings shed light on how virtual groups and teams can enhance their collaboration by leveraging the power of communication channels. By understanding the impact of non-verbal cues and signals, organizations can foster stronger connections and achieve more successful outcomes.

TITEL: Nature Rights to Promote Sustainability: A Question of Voice or Compensation for Ecosystems?

AUTOR*INNEN: Stephan Geschwind, Johann Graf Lambsdorff, Deborah Voß

ABSTRACT: In response to the climate crisis and the accelerated destruction of ecosystems, the idea of giving nature legal personality has gained increasing interest. Such legal status endows ecosystems with the right to exist, evolve, and thrive without interference from human-caused harms like pollution and resource depletion. They are no longer treated as property but have legal standing in court and can sue for damage compensation when harmed. The question whether such rights are sufficient to promote sustainable behavior among resource users has not clearly been answered yet. Arguably, both nature's ability to stand in court and to receive damage compensation can promote sustainable behavior. We propose a novel experimental design that disentangles the two possible mechanisms. We combine a dynamic multi-round common-pool resource game with a novel third-party one-shot damage compensation scheme. We vary the specifics of the third-party damage compensation scheme in a 2x2 factorial design: One dimension varies the standing of nature in court as an own legal entity and a pure commodity. The third-party is either instructed to take into consideration the state of the pool or the efficiency of resource use. The task is incentivized based on an AI's assessment of role fulfillment. The other dimension varies the allocation of damage compensation. In one scenario, it is an anthropocentric transfer between participants. In the other scenario, it is an ecocentric restitution of the pool.

Session 2 (Donnerstag 26.09., 09:00 Uhr)

S11	S12	S14	S15	S16
Ethics and Fairness I	AI and Strategic Interaction II	AI and Language	Cooperation and Prosociality II (Sessionsprache: Deutsch)	Effort & Performance
Redistribution, Moral Hazard, and Voting by Feet: An Experiment Alisa Frey (Heinrich-Heine-Universität Duesseldorf)	ChatGPT's financial discrimination between rich and poor – misaligned with human behavior and expectations. Dmitri Bershadskyy (Otto-von-Guericke Universität)	Do You Trust Your Own Voice? Jannik Greif (Otto-von-Guericke Universität Magdeburg)	Payoffs, Beliefs, and Cooperation in Infinitely Repeated Games Maximilian Andres (Universität Potsdam)	I Know What You Did Last Summer: How Past Performance Affects Future Performance Evaluation Fabian Felser (Technische Universität Dresden)
Populist propaganda and anti-migration stances: An experimental investigation Christian Koch (Universität Wien)	A Behavioral Taxonomy of 2x2 Games Rosemarie Nagel (ICREA-UPF-BSE)		Why do People Follow an Example? Maximilian Kuntze (Universität Vechta)	The Effect of Task (Mis)Matching and Self-Selection on Intrinsic Motivation and Performance Louis Strang (Universität zu Köln)
Corporate Responsibility Increases Consumers' Willingness to Pay: Evidence from two Framed Field Experiments Julian Conrads (Technische Hochschule Mittelhessen)	Karma: An experimental investigation of a 'closed' intertemporal token system Heinrich Nax (ETH/UZH)	Seeing is believing: How (X)AI-augmented advice guides human task-solving behavior Florian E. Sachs (Universität zu Köln)	Motivational Effects of Monetary Incentives and Social Recognition on Parental Engagement Nicole Middendorf (Universität Vechta)	

Ethics and Fairness I

TITEL: Redistribution, Moral Hazard, and Voting by Feet: An Experiment

AUTOR*INNEN: Alisa Frey

ABSTRACT: In a laboratory experiment, participants first choose between two purely redistributive tax regimes and then generate income through a real effort task. High taxes insure against unfavorable gross incomes, but increase effort-reducing moral hazard. The results suggest that subjects' initial tax choices behind a veil of ignorance are quite heterogeneous. Overconfident (underconfident) subjects are more (less) likely to choose low taxes. When familiar with the task and their relative gross incomes, subjects choose the tax rate opportunistically: High-income subjects tend to choose the low rate, while low-income subjects tend to choose the high rate myopically. Voting by feet eventually leads to a high proportion of subjects choosing the low tax. A control treatment in which the tax rate is exogenous controls for selection and helps to identify the moral hazard effect, which is mainly driven by those above the median. A further treatment in which income is randomly given helps to quantify subjects' redistributive preferences arising from income uncertainty.

TITEL: Populist propaganda and anti-migration stances: An experimental investigation

AUTOR*INNEN: Christian Koch, Jean-Robert Tyran

ABSTRACT: Populism and anti-migration sentiments have gained prominence around the world in recent years. Our study probes the scope and limits of political framing in shaping attitudes towards immigration, employing an experimental design. We explore the effectiveness of different populist messages across different economic conditions through framing manipulations. Individuals' inclinations toward pro- or anti-migration stances hinge on both the economic ramifications of migration and their psychological construal of it. For the latter, we focus on two common themes in populist rhetoric: portraying immigrants as free riders and framing opposition to immigration as the 'right thing to do' to safeguard fellow citizens. We find that anti-migration framing can be highly effective, even in our sample of university students, doubling opposition to immigration. Intriguingly, this effect persists even when immigration lacks a direct monetary impact on individuals. Notably, individuals who perceive real-world immigrants as free riders are more inclined to oppose migration in our experimental settings. These results underscore the potential influence of political messaging in shaping public perceptions of immigration.

TITEL: Corporate Responsibility Increases Consumers' Willingness to Pay: Evidence from two Framed Field Experiments

AUTOR*INNEN: Julian Conrads, Alexandra Eyberg, Bernd Irlenbusch, Maivand Sarin

ABSTRACT: Do consumers reward companies for engaging in corporate responsibility (CR)? This longstanding question is vividly discussed, at least since Milton Friedman's (1970) famous doctrine: "The social responsibility of business is to increase its profits." Clear evidence on the question is difficult to obtain in naturally occurring purchasing

environments because of many confounding effects like brand reputation, product quality, and product appearance. CR predominantly focuses on two dimensions, i.e., (i) the social dimension, like providing fair working conditions, and (ii) the ecological dimension, like using recycled materials in production. Through a framed field experiment in backpack sales for young adults, we provide new evidence on this debate. We find that CR information about both dimensions combined increases consumers' willingness to pay (WTP) by more than 17 percent. In a separate study, we ask which of the two is more effective and compare the effects of CR information in the social dimension with CR information in the ecological dimension. Interestingly, we find no significant difference in consumers' WTP between these two dimensions, suggesting that it might not matter what kind of CR activity a company pursues as long as it does something.

AI and Strategic Interaction II

TITEL: ChatGPT's financial discrimination between rich and poor – misaligned with human behavior and expectations.

AUTOR*INNEN: Dmitri Bershadskyy, Florian E. Sachs, Joachim Weimann

ABSTRACT: ChatGPT disrupted the application of machine-learning methods and drastically reduced the usage barrier. Chatbots are now widely used in a lot of different situations. They provide advice, assist in writing source code, or assess and summarize information from various sources. However, their scope is not only limited to aiding humans they can also be used to take on tasks like negotiating or bargaining. To understand the implications of Chatbot usage on bargaining situations, we conduct a laboratory experiment with the ultimatum game. In the ultimatum game, two human players interact: The receiver decides on accepting or rejecting a monetary offer from the proposer. To shed light on the new bargaining situation, we let ChatGPT provide an offer to a human player. In the novel design, we vary the wealth of the receivers. Our results indicate that humans have the same beliefs about other humans and chatbots. However, our results contradict these beliefs in an important point: Humans favor poor receivers as correctly anticipated by the humans, but ChatGPT favors rich receivers which the humans did not expect to happen. These results imply that ChatGPT's answers are not aligned with those of humans and that humans do not anticipate this difference.

TITEL: A Behavioral Taxonomy of 2x2 Games

AUTOR*INNEN: Amil Camilo, Fabrizio Germano, Rosemarie Nagel

ABSTRACT: This paper presents a comprehensive analysis of the 2x2 "ordinal" games to derive a taxonomy of games, in line with behavioral rules and experimental evidence. We introduce a graph-theoretic method for computing similarity classes based on the neighborhood structure of games and the discontinuities of behavior within the games. Using this method, we derive taxonomies based on standard game theory and behavioral rules, such as iterated best responses to uniform random play (level-k) and rules based on social preferences (e.g., the equal split or social maximum). We contrast these taxonomies with actual behavior to derive a behavioral rule and experimental data-based taxonomy of 2x2 games. Our approach should

provide a useful tool for experimental design and algorithmic game theory and identify strategically interesting games and classes of games that have received little attention in the economic literature.

TITEL: Karma: An experimental investigation of a 'closed' intertemporal token system

AUTOR*INNEN: E Elokda, H Nax, S Bolognani, F Dörfler

ABSTRACT: A system of non-tradable credits that flow between individuals like karma, hence proposed under that name, is a mechanism for repeated resource allocation that comes with attractive efficiency and fairness properties, in theory. In this study, we test karma in an online experiment in which human subjects repeatedly compete for a resource with time-varying and stochastic individual preferences or urgency to acquire the resource. We confirm that karma has significant and sustained welfare benefits even in a population with no prior training. We identify mechanism usage in contexts with sporadic high urgency, more so than with frequent moderate urgency, and implemented as an easy (binary) karma bidding scheme as particularly effective for welfare improvements: relatively larger aggregate efficiency gains are realized that are (almost) Pareto superior. These findings provide guidance for further testing and for future implementation plans of such mechanisms in the real world.

AI and Language

TITEL: Do You Trust Your Own Voice?

AUTOR*INNEN: Jannik Greif

ABSTRACT: Deepfakes, highly realistic manipulations of audio-visual content, have reached a level of sophistication that makes it challenging for individuals to distinguish between authentic and fabricated material. Deepfakes are becoming increasingly prevalent and utilized in various contexts. They can alter perceptions of reality, potentially leading to misinformation and manipulated narratives. In the context of audio communication, deepfakes involve modifications such as changing the pitch, tone, or cloning voices. This can create persuasive interactions, influencing trust levels depending on whether the deepfake is detected. Prior research indicates that people struggle to differentiate between authentic and deepfake audio content. This study focuses on the strategic use of audio deepfakes, examining whether individuals are inclined to use audio deepfakes when given the option and whether these are employed to build initial trust for reciprocation or exploitation. As the tools for creating deepfakes are constantly improving and less technical expertise is required, it remains unclear to what extent this will be used in day-to-day communication and what the consequences might be. To investigate this, a lab experiment is conducted, utilizing a trust game with one-sided authentic and deepfake audio communication. This study contributes to our understanding of deepfake usage and its implications in the context of digital audio communication.

TITEL: Seeing is believing: How (X)AI-augmented advice guides human task-solving behavior

AUTOR*INNEN: Dmitri Bershadskyy, Kai Heinrich, **Florian E. Sachs**

ABSTRACT: Artificial Intelligence (AI) based advice can affect human decisions during task-solving processes in various ways. A particular concern for tasks with high repetition rates, such as anomaly detection in quality assurance, is the risk of users accepting AI advice without further consideration – to save time. This study investigates this issue through a laboratory experiment utilizing eye-tracking technology, focusing specifically on anomaly detection. We examine two types of AI advice: one that merely predicts product quality and another that, in addition to the prediction, provides an AI-generated heatmap highlighting detected anomalies' locations. Our research addresses three key questions: (i) Does anomaly detection accuracy improve with AI recommendations? (ii) Does the time participants take to make decisions change when AI advice is utilized? (iii) Do participants focus on the areas highlighted by AI-generated heatmaps? The results indicate that AI recommendations lead to higher accuracy in anomaly detection and that participants spend more time making decisions when receiving AI advice. Eye-tracking data corroborates these findings. These results underscore the potential benefits and challenges of integrating AI into repetitive decision-making tasks, with implications for both the design of AI systems and user training.

Cooperation & Prosociality II

TITEL: Payoffs, Beliefs, and Cooperation in Infinitely Repeated Games

AUTOR*INNEN: Maximilian Andres, Lisa Bruttel, Juri Nithammer

ABSTRACT: This paper studies the interaction of beliefs, payoff parameters, and cooperation in the infinitely repeated prisoner's dilemma. We show, formally and in laboratory experiments, that a player's belief about the probability of cooperation by their opponent moderates the effect of changes in the payoff parameters on cooperation. If beliefs are high, increasing the gain from unilateral defection has a large negative effect on cooperation, while increasing the loss from unilateral cooperation has a negligible negative effect. Conversely, if beliefs are low, this relationship is reversed: increasing the gain has a smaller negative effect, while increasing the loss has a large negative effect on cooperation.

TITEL: Why do People Follow an Example?

AUTOR*INNEN: Gerald Eisenkopf, **Maximilian Alex Kuntze**

ABSTRACT: Recent meta-analyses on leading by example indicate a strong influence of decisions by a first-moving leader on decisions by second-moving followers in social dilemma situations. However, in the standard leading-by-example game, outcome and intention motives of followers' reciprocation are aligned, making their distinct effects difficult to assess. Using a public goods game, we disentangle the relevance of a leader's intention and two outcome measures—her impact and her cost—by randomizing impact and cost after the leader's contribution decision but before the followers observe it and decide themselves. We

also investigate how information about these measures affects long-term cooperation. Preliminary results indicate that the leader's intentions and actual impact are the most significant motives for followers' reciprocity, while the leader's costs matter, but not as strongly. Furthermore, our findings suggest that random distortions of costs and contributions do not undermine long-term cooperation.

TITEL: Motivational Effects of Monetary Incentives and Social Recognition on Parental Engagement

AUTOR*INNEN: Vanessa Mertins, Nicole Middendorf

ABSTRACT: In our study, we analyze the motivational impact of a monetary incentive versus social recognition compared to a control group without incentives on a heterogeneous group of parents who are tasked with supporting their children in reading activities. Several hundred 6-8-year-old elementary school students between first and second grade practice their reading skills in a paired loud reading procedure within a randomized field experiment, with the help of their parents. By providing this method and selected daily reading texts, we equip parents with a tool that can directly enhance the children's reading performance and indirectly increases their motivation to read. To fundamentally encourage parents to invest time in the parent-child tandem, the effects of a monetary incentive versus social recognition compared to a control group without incentives are tested.

Effort & Performance

TITEL: I Know What You Did Last Summer: How Past Performance Affects Future Performance Evaluation

AUTOR*INNEN: Philipp Richter, Peter Schäfer, Fabian Felser

ABSTRACT: We address a critical challenge in management accounting: mitigating cognitive bias in subjective performance evaluation. Do supervisors unconsciously rely on past evaluations when rating their employees' performance? This can harm employees' perceptions of fairness and motivation, and psychology research indeed provides evidence of such a bias. We explore the cognitive mechanisms underlying this bias and evaluate the effectiveness of different strategies to mitigate it through three experiments. First, we investigate whether supervisors who are made aware of the bias in a pre-evaluation stage or are required to justify their judgments in a post-evaluation stage put less weight on irrelevant past performance. Second, we investigate whether framing of evaluation moderates the decision process in taking past evaluations into account. Third, we investigate whether supervisors are less prone to the bias if they are more familiar with their team members' task. Our project will have significant implications for designing effective incentive systems. Many firms provide supervisors with discretion in performance evaluation, but biases in supervisors' ratings render subjective performance evaluations ineffective. It is thus crucial for management accounting professionals to understand these biases and have techniques to reduce them.

TITEL: The Effect of Task (Mis)Matching and Self-Selection on Intrinsic Motivation and Performance

AUTOR*INNEN: Timo Freyer, Jonas Radbruch, Sebastian Schaub, Louis Strang

ABSTRACT: This paper investigates how the self-selection of tasks affects worker performance. Specifically, it investigates the impact of aligning tasks with workers' preferences and the effect of providing workers with greater autonomy in choosing their tasks. To answer these questions, we conducted an online experiment in which participants engaged in one of two real-effort tasks. We exogenously varied whether participants were either assigned their preferred or non-preferred task, or if they had the opportunity to actively self-select their task. The results show that participants who were randomly assigned their preferred task or self-selected a task increased their output by about 23%–36% of a standard deviation compared to those who were assigned their non-preferred task. This increase in output is linked to both enhanced productivity and extended time spent working on the task. In essence, our results underscore that workers' performance depends crucially on whether they work on their preferred task. Importantly, our results also document that granting workers decision autonomy in task selection reinforces the performance increase.

Session 3 (Donnerstag 26.09., 11:00 Uhr)

S11	S12	S14	S15	S16
Cheating and Honesty	AI and Humans I	AI and Preferences	Trust	Digitalization (Sessionsprache: Deutsch)
Prevalence and perceptions of undeclared work and tax evasion Lilith Burgstaller (Universität Freiburg)	Delegation to Pricing Algorithms: Experimental Evidence Nina Rulié (Heinrich-Heine-Universität Düsseldorf)	Learning to be kind: neural networks as reciprocal altruists in allocation games Christian Feige (unabhängiger Forscher)	Do people trust Artificial Intelligence? A Repeated Trust Game Experiment Dinithi Jayasekara (Singapore University of Technology and Design/ Lee Kuan Yew)	
Ethical Risks of Delegation to AI Nils Köbis (Universität Duisburg-Essen)	Searching for Stars or Avoiding Catastrophes: The Role of Loss Aversion for Trust in AI Julia Werner (Universität Konstanz)	Do Personalized AI Predictions Change Subsequent Decision-Outcomes? Artificial versus Swarm Intelligence Paul M. Gorny (Karlsruher Institut für Technologie)	Improving Trust as a Tool for Improving Tax Compliance Matthias Kasper (Walter Eucken Institut)	Bridging the Communication Gap: Real-Time Visualization of Eye Contact to Support Cooperation in Virtual Teams Tom Reuscher (Karlsruher Institut für Technologie)
Beliefs and Group Dishonesty: The Role of Strategic Interaction and Responsibility Rainer Michael Rilke (WHU - Otto Beisheim School of Management)	AI or human? Applicants' decisions in discrimination settings Luisa Santiago Wolf (Universität zu Köln)	Don't do what I might choose, but please do what will be better: An algorithm's role, feedback, & algorithm reliance Irenaeus Wolff (Universität Konstanz)	Handshakes and Trust in Virtual Reality Lucas Braun (RWTH Aachen)	How digital media markets amplify news sentiment Lara Berger (Universität zu Köln)

Cheating and Honesty

TITEL: Prevalence and perceptions of undeclared work and tax evasion

AUTOR*INNEN: Lilith Burgstaller, Lars Feld, Katharina Pfeil

ABSTRACT: Undeclared work in the household is prevalent in all OECD countries (OECD, 2021). In these transactions, seller and buyer actively collude to not pay taxes, social security contributions as well as other insurances. While policy makers debate the extent and measures against tax evasion, the true extent of it is hard to measure as it is usually concealed from authorities and offenders refrain from admitting to it when asked. In this study we show that collaborative tax evasion is prevalent: In a list experiment, some 18% of our sample indicate to have experience with undeclared work in the household. When asked directly this share amounts to 20% of individuals. At the same time, 52% of individuals say they know someone in their social circle who has experience with undeclared work in the household.

Regarding perceptions of collaborative tax evasion in the household, we show that respondents also expect collaborative tax evasion in the household to be quite common: In the mean, they expect that 43% of households in Germany have experience with undeclared work in the household. Moreover, on average respondents think that 58% of households in Germany think that is justifiable to agree on undeclared work. Heterogeneity analyses show that these perceptions vary by gender, age and income. Moreover, these perceptions suggest that collaborative tax evasion is considered much more common than tax evasion by the self-employed.

TITEL: Ethical Risks of Delegation to AI

AUTOR*INNEN: Iyad Rahwan, Zoe Rahwan, Jean-Francois Bonnefon, Tamer Ajaj, Clara Bersch, Nils Köbis

ABSTRACT: Rationale: A growing exists trend to delegate tasks to algorithms, for example smart pricing options on platforms like Airbnb. Yet, algorithms that autonomously act on people's behalf might break legal or ethical rules, even without humans being aware of it, such as the recent evidence of pricing algorithms colluding autonomously. This raises a fundamental question about Artificial Intelligence safety: can the ability to delegate tasks to machines increase human engagement in unethical behavior? Methods: To examine this question across different types of algorithmic settings we conducted four pre-registered, large-scale online experiments (total N= 3217). As a measure for ethical behavior we used the die-rolling task, a well-established paradigm where participants are instructed to report the observed outcome of a private die roll but get paid according to the number they report. In all experiments we compare the degree of dishonesty in different delegation settings to a baseline where people report the die rolls themselves. Results. First, our results show that individuals are more prone to unethical behavior when delegating tasks to machines than when performing these tasks themselves, with only 5-10% behaving unethically in the latter scenario. Second, we show that the manner in which humans program the machines—e.g., using rules, training data, or high-level goals—can qualitatively alter the temptation towards unethical behavior, causing as many as 85% of participants to behave unethically in some conditions. Together, the quantitative increase in opportunities to delegate, coupled with the qualitative change in temptation, combine to

substantially increase unethical behavior. However, we also identified two effective strategies to mitigate this risk: allowing participants to choose whether to delegate to a machine or undertake the task themselves significantly increased honest behavior, with a preference for self-engagement rising to 74% after experience with both human and machine delegation. Additionally, using natural language for delegation, a feature now common in modern chatbot technology, notably reduced unethical delegation. Conclusion: These results highlight the contexts under which risk of increased unethical behavior arises when delegating to AI and underscore the importance of considering human factors in AI safety.

TITEL: Beliefs and Group Dishonesty: The Role of Strategic Interaction and Responsibility

AUTOR*INNEN: Fabio Galeotti, Rainer Michael Rilke, Eugenio Verrina

ABSTRACT: Dishonest behavior often occurs in groups where actions are interconnected and beliefs about others' behavior may play an important role. We study the relationship between beliefs and dishonesty, focusing on the impact of the nature of the strategic interaction (complements or substitutes) and the reduced feeling of responsibility that arises from acting together with other group members. In settings of strategic complements, we observe that individuals tend to lie more, the more they believe their counterpart to be dishonest. Conversely, in settings of strategic substitutes, individuals tend to lie less as their belief about their counterpart's dishonesty increases. Acting together instead of acting alone --- while holding incentives and beliefs constant --- does not influence the relationship between beliefs and behavior in strategic complements. However, individuals with higher lying costs lie less in strategic substitutes when they are the only active member of the group. Our findings suggest that both beliefs and the type of strategic interaction strongly shape group dishonesty, while responsibility plays a minor role.

AI and Humans I

TITEL: Delegation to Pricing Algorithms: Experimental Evidence

AUTOR*INNEN: Hans-Theo Normann, Nina Rulié, Olaf Stypa, Tobias Werner

ABSTRACT: In a market experiment, we analyze the propensity of participants to delegate their decisions to an algorithm. The algorithm used to make pricing decisions in the experiment results from extensive (offline) simulations with a self-learning reinforcement learning algorithm (Q-learning), which has been shown to be able to collude on noncompetitive prices tacitly and is more collusive than humans in duopolies. A participant who wants to delegate her pricing decisions to the algorithm has an incentive to do so since it can facilitate market coordination and increase profits. We tell subjects that "the algorithm is free to use and aims to maximize your revenue over the course of the experiment.". In a Bertrand duopoly with inelastic demand, we analyze three treatments. In the baseline treatment, two human players play without any algorithm. In Outsourcing, participants can delegate their decisions to the algorithm and are fully committed, as they cannot override the algorithm's decisions. The

Recommendation treatment is like Outsourcing, but subjects can overwrite the algorithm's decision. Our results show that subjects delegate significantly more in Recommendation than in Outsourcing. Prices are lower in Recommendation compared to Outsourcing and Baseline. Prices are not significantly different between Outsourcing and Baseline. We also compare the price level for different subgames (human-human vs. human-algorithm).

TITEL: Searching for Stars or Avoiding Catastrophes: The Role of Loss Aversion for Trust in AI

AUTOR*INNEN: Lara Suraci, Julia Werner, Irenaeus Wolff

ABSTRACT: Artificial intelligence (AI) is rapidly advancing and becoming ubiquitous, taking over tasks once reserved for humans and even surpassing human performance in some areas. Despite these advancements, the level of reliance on AI is still often suboptimal, frequently falling short of the potential benefits that could be realized through its utilization. Research on individuals' willingness to rely on AI for decisions has been growing, but results remain largely inconclusive regarding which factors influence this reliance and how. We examine how situational factors, like the goal of the choice (choosing the best option [stars] vs. avoiding particularly bad options [catastrophes]) and decision-authority (advice vs. delegation), as well as loss aversion as an individual factor, affect the reliance on AI. To investigate this, we conduct an experiment where subjects can, depending on the treatment, receive advice or delegate a decision to one of three different agents (AI-agent, human-agent, or human-agent with AI-agent support) to assist them with a complex investment task. Additionally, they have the option to decide independently. We expect that as the perceived risk of a negative outcome increases, which can be influenced by both loss aversion and situational factors (i.e., stars vs. catastrophes), people will be less likely to delegate, especially to an AI. However, this pattern does not hold for advice. Here, we think that people prefer to receive advice from an AI-agent compared to a human-agent, regardless of the perceived risk of a negative outcome.

TITEL: AI or human? Applicants' decisions in discrimination settings

AUTOR*INNEN: Paula Thevißen, David Stommel, Luisa Santiago Wolf

ABSTRACT: While there is agreement on the great potential for efficiencies and savings that the use of artificial intelligence (AI) can bring to recruitment, there is also an ongoing debate about the ethical and legal implications of hiring algorithms. The existing literature on the perception of hiring algorithms is ambiguous, especially for settings that incorporate the risk of discrimination. In the paper, we study whether anticipated discrimination influences applicants' preferences for a hiring algorithm. We conduct an online experiment based on the design of Dagnies et al. (2022) in which applicants are asked to decide whether a human manager or an AI should take their hiring decision. The novelty in our approach is the simulation of discriminatory settings. We distinguish between taste-based and statistical discrimination and address a wide range of characteristics on which grounds applicants can be discriminated against. This extends our insights to a broader group of applicants than usually studied in gender discrimination settings. Our results show that the type of discrimination influences applicants' preference for an AI or human recruiter. In the taste-based discrimination setting, in which the median applicant expects to be discriminated against, we find a preference to be hired by an AI over a human. Suggestive evidence indicates that applicants consider the AI to be less biased than a human. In the statistical discrimination

setting, the median applicant expects the AI to better understand information on performance differences related to statistical characteristics of the applicants. Accordingly, the median applicant from a statistically disadvantaged group prefers a human manager over AI, while members of statistically advantaged groups prefer an AI. Finally, our results suggest that announcing the use of AI in the recruitment process has the potential to affect the composition and thus the diversity of the applicant pool.

AI and Preferences

TITEL: Learning to be kind: neural networks as reciprocal altruists in allocation games

AUTOR*INNEN: Christian Feige

ABSTRACT: The aim of this project is to train a neural network based on reinforcement learning to act as a reciprocal altruist. This means that, in order to maximize its own payoffs, the neural network learns to treat other agents kindly by showing empathy with their needs. The experiment evaluates the performance of the neural network in a repeated interaction with another agent in an allocation game. In this game two players make voluntary contributions to a common project, which pays a reward if total contributions exceed a threshold value. The neural network takes the role of a less productive player who benefits from convincing its more productive partner to carry most or all of the contribution burden. The other player is a computer agent with context-dependent other-regarding preferences ranging from altruism via inequality aversion to spite. Similar to intrinsic reciprocity, this player's preferences change in reaction to advantageous or disadvantageous payoff inequality. Different treatment conditions vary the accuracy of signals about the partner's current preferences, the partner's aversion to disadvantageous payoff inequality, and the reward for reaching the threshold value. Preliminary results show that reliable information about partner preferences generally helps the neural network to significantly increase its own payoff, independently of the size of the reward or the partner's extent of inequality aversion. The only exception is the combination of high incentives and an inequality-averse partner, in which case the reciprocal altruist appears to earn the highest payoffs for moderate levels of uncertainty.

TITEL: Do Personalized AI Predictions Change Subsequent Decision-Outcomes? Artificial versus Swarm Intelligence

AUTOR*INNEN: Paul M. Gorny, Christina Strobel

ABSTRACT: Artificial Intelligence (AI) substantially improves our ability to predict future choices. Decision-makers are therefore likely to get in touch with ever more personalized predictions about their own decisions in the near future. We investigate how salient predictions about future choices affect these choices with a particular focus on predictions made by (i) an AI or (ii) other humans. In a first experiment, humans generate predictions of other humans' choices in one-shot dictator games using an incentive-compatible mechanism. The AI is then calibrated to match the humans' accuracy in making predictions. In a second experiment, participants make a decision in such a dictator game either without a prediction (baseline) or with a prediction from the AI or the humans, depending on treatment. Beyond comparing

sharing behavior across these three treatments, we also consider how predictions differ from presenting participants with statistical information to net out any effects of descriptive norms affecting behavior. We discuss our results together with potential mechanisms. Our project contributes to understanding the role of technocratization, i.e., the rising role of technology or technical solutions when it comes to future choices.

TITEL: Don't do what I might choose, but please do what will be better: An algorithm's role, feedback, & algorithm reliance

AUTOR*INNEN: Lara Suraci, Robin Cubitt, Chris Starmer, Irenaeus Wolff

ABSTRACT: In the literature on the determinants of algorithm reliance, large gaps remain regarding environments where preferences matter. Our experiment aims to address the algorithm's role in the realm of decision-making under risk, and adds additional insights with respect to the effect of dissatisfying feedback. We observe that the majority of our participants chooses against an algorithm that tries to mimic their preferences while avoiding dominated options. In contrast, a majority of the participants (initially) choose an algorithm that maximises expected value. Over time, algorithm reliance decreases in all treatments. However, this trend seems more pronounced in treatments that include outcome feedback. Interestingly, participants facing the preference-mimicking algorithm react to outcome feedback more strongly than those facing the expected-value maximising algorithm. We interpret the findings such that people feel they know what they would want to choose, but they do not know whether the preference-mimicking algorithm is able to reproduce that behaviour---potentially adding errors---but when they have the chance to choose (or 'commit to') a strategy that is 'better overall', they are happy to choose that. Data from our post-experimental questionnaire further supports the interpretation.

Trust

TITEL: Do people trust Artificial Intelligence? A Repeated Trust Game Experiment

AUTOR*INNEN: Dinithi N. Jayasekara, Benjamin Prisse, Deng Ruotong, Jun Quan Ho

ABSTRACT: Uncovering how human-artifice interactions manifest trust is important for the effective deployment of Artificial Intelligence (AI) technologies. In this paper, we investigate public trust in AI through a repeated trust experiment conducted online. Participants engage in a 20-period trust game with an AI assuming the roles of the Truster and Trustee, respectively. The Trustee begins with an endowment of 10, while the AI starts with an endowment of 0. All periods are independent. In each period, there is a 10% probability that the AI does not respond due to an error. This experimental design simulates a scenario in which participants must decide whether to invest in a firm producing new AI technologies. In Part 2, participants play the same game in reverse positions to examine their reciprocity towards AI. In Part 3, the AI starts with an endowment of 10, simulating an interaction with an established firm producing AI technologies. Additionally, participants completed a Dictator Game, an Ultimatum Game, a sociodemographic survey, and an AI questionnaire designed specifically for this experiment to measure their knowledge and beliefs regarding AI. Our experimental design accounts for three featured AI systems,

namely, Social AI, Exploratory AI, and Competitive AI. They mimic different types of personalities" that imitate real-world functionalities of AIs (e.g., assistant, competitor, learner). Social AI mimics AI systems that consistently reward individuals, for example, chatbots and generative AI tools. Exploratory AI imitates AI systems that are uncertain, ambiguous, risky, and driven by curiosity and a desire to uncover hidden patterns or knowledge within data. In contrast, Competitive AI is always strategic as those displayed by competitive games such as AlphaGo. We include additional treatments, Learning AI, to mimic AI systems that learn by observing the interactions with individuals, and two control treatments that directly compare interactions with humans and AI (Human-AI, Human-Human). We differentiate the trust decisions of individuals based on the AI they engage with. Preliminary results indicate that subjects gradually increase their trust in AI, and initial beliefs are persistent. Subjects with initially low trust show the most significant increase. We also observe that subjects adjust their trust based on the expected benefits of the treatment. These findings suggest that interactions with AI are logically driven, and individuals can quickly be convinced to use AI if they perceive direct benefits."

TITEL: Improving Trust as a Tool for Improving Tax Compliance

AUTOR*INNEN: Matthias Kasper, James Alm

ABSTRACT: It is increasingly recognized that individual trust in government institutions affects the willingness to pay taxes. However, specific ways in which this trust can be changed remain largely unexamined and unexplained. This paper analyzes the effects of trust on tax compliance using experimental methods. We first establish a basic fiscal system in which taxes are used to finance public expenditures but where individuals are able to evade their taxes with some probability of detection and punishment. We then systematically alter various features of the fiscal system in ways that affect trust in the government. In particular, we introduce systematic variation in tax system design such as voting on tax rates, fairness (procedural, distributive, retributive), and social norms. By comparing trust and compliance across conditions, we are able to determine how the policy change affects trust and so improves tax compliance.

TITEL: Handshakes and Trust in Virtual Reality

AUTOR*INNEN: Lucas Braun, Özgür Gürerk, Bernd Irlenbusch

ABSTRACT: We investigate whether a virtual handshake, simulated using haptic feedback technology, can increase trust in business relationships that take place in virtual environments. We focus on the use of vision-based hand tracking in combination with haptic feedback gloves to provide a sense of touch and physical interaction in virtual reality, and its potential to enhance trust and collaboration between individuals. Through a series of experiments, conducted using virtual reality headsets, we examine the effects of virtual handshakes on trust in a variant of the Moonlighting Game (Abbink et al. 2000), where a handshake is used to seal the non-binding contract. The results of the study are expected to provide insight into the potential of haptic feedback technology to improve trust and collaboration and have implications for the future of virtual business interactions.

Digitalization

TITEL: Algorithmusaversion vs. Verlustaversion **CANCELLED!**

AUTOR*INNEN: Ibrahim Filiz, **Florian Kirchhoff**, Thomas Nahmer, Markus Spiwoks

ABSTRACT: Algorithmusaversion beschreibt eine Verhaltensanomalie, nach der Menschen effizienteren, algorithmusbasierten Systemen misstrauen und stattdessen menschliches Urteilsvermögen bevorzugen. Wirtschaftssubjekte laufen damit Gefahr, nicht ihren maximal erreichbaren Nutzen zu realisieren. Diese Studie soll einen Beitrag zu der Frage leisten, wie Algorithmusaversion reduziert werden kann. Im Rahmen eines Laborexperiments wird dafür überprüft, ob die bereits intensiv erforschte, wirkungsvolle Verhaltensanomalie der Verlustaversion zur Reduktion von Algorithmusaversion beitragen kann. Tatsächlich zeigt sich, dass das Gegenteil der Fall zu sein scheint: Die Bereitschaft, einen im Vergleich zu einem menschlichen Experten erkennbar leistungsfähigeren Algorithmus einzusetzen, geht sogar zurück, wenn bei der Entscheidung ein Verlust droht. Dieser Befund stützt andere Forschungsergebnisse, wonach Algorithmusaversion bei schwerwiegenden möglichen Konsequenzen verstärkt auftritt. Zur Verbreitung algorithmusbasierter Systeme scheint es demnach angebracht zu sein, die mit dem Einsatz verbundenen Chancen durch ein positives Framing zu betonen und zu bewerben.

TITEL: Bridging the Communication Gap: Real-Time Visualization of Eye Contact to Support Cooperation in Virtual Teams

AUTOR*INNEN: Petra Nieken, **Tom Frank Reuscher**

ABSTRACT: In recent years, managers face the challenge of establishing flexible work models that meet their employees' growing desire to be physically independent from a single fixed workspace, while still ensuring productivity. To accomplish this, it is necessary to guarantee an efficient exchange of information among team members situated in distant locations. In most organizations, interaction via video-meetings has become the default for remote cooperation, as its visual component provides the richest form of computer-mediated communication. In this study, we investigate the potential benefits of an innovative video-meeting system that bridges the gap between computer-mediated and face-to-face communication by adaptively visualizing episodes of eye contact between team members. Specifically, we conduct a controlled lab experiment to test whether the ability to initiate and perceive bidirectional eye contact in remote interaction effects cooperation and coordination in an iterated Weakest Link Game. Given that the coordination failure is due to strategic uncertainty, we designed two treatments that allow players to interact in a pre-play (not game-related) video-meeting for eight minutes. In the Baseline Treatment, we use a standard video-meeting system that enables players to see and hear each other. In the Eye Contact Treatment, we use an adaptive video-meeting system that also processes gaze information. In particular, we inform players if they are looking at each others video feeds simultaneously. Our research aims to provide valuable insights into the ongoing debate regarding the optimal communication methods in professional settings and the role of eye contact and gaze patterns in team cooperation. By assessing the

potential of enhanced video-meeting systems to facilitate better cooperation and coordination, we seek to contribute to the development of more effective remote work practices.

TITEL: How digital media markets amplify news sentiment

AUTOR*INNEN: Lara Marie Berger

ABSTRACT: Capturing attention through appealing headlines is far more important for news companies in digital than in analogue media markets, but it is so far unclear if this changes news content. This paper provides evidence that this shift in incentives enhances the sentimental slant of news headlines. A comparison of online and offline versions of the same newspapers illustrates that headlines online are more often formulated emotionally. An experiment with professional journalists reveals that this difference can be at least partially explained by an increased incentive to generate clicks: If journalists are compensated relative to the click-rates their headlines receive, they significantly more often put headlines containing emotional words on top of a given article. A second experiment shows that such an amplification of sensationalist framing has economic implications in the short run, as emotional headlines can translate into emotional reactions and distortions in expectations of their readers.

Session 4 (Donnerstag 26.09., 14:00 Uhr)

S11	S12	S14	S15	S16
Volunteering and Charity (Sessionsprache: Deutsch)	AI and Humans II	Bargaining	Social Norms and Reference Points	Health
Recruiting Episodic Volunteers: A Field Experiment on Gift-Giving vs. Requesting a Favor Maximilian Hiller (Universität Vechta)	Can the Use of Large Language Models in Higher Education Empower Students to Increase their Learning Success? Eduard Buzila (Otto-von-Guericke Universität Magdeburg)	The Economic Effects of Remote-Bargaining vs. Face-to-Face Bargaining Holger Rau (Georg-August Universität Göttingen)	Sanctions and Social Norms Amanda März (Walter Eucken Institut)	Human Infection Risk Assessment in Social Networks: The Role of Network Characteristics Martin Strobel (Maastricht University)
Boosting Performance Quality in Episodic Micro-Volunteering: A Collaborative Field Experiment. Bernd Josef Leisen (Universität Vechta)	Appreciation to biased algorithms, and aversion from unbiased algorithms. Yefim Roth (Ori Plonsky) (University of Haifa and Technion - Israel Institute of Technology)	Between egalitarianism and utilitarianism: experimental evidence in unstructured bargaining Behnud Mir Djawadi (Universität Paderborn)	Social Networks and Coordination: A Simulation Experiment Blanca Tena Estrada (Universität Kassel)	An experiment on decision support in medicine Anna Katharina Stirner (Universität zu Köln)
	Overreliance, Backfire, and the Importance of (Mis)trust in Human-Automation Interactions Yefim Roth (University of Haifa)	Lie Against AI: Revealing Private Information through AI in an Economic Experiment Dmitri Bershadskyy (Nina Ostermaier) (Otto-von-Guericke Universität Magdeburg)	Determinants and Dynamics of Collaborative Exploitation Till Vater (Maastricht University)	Sow the Doubt, Reap the Inefficiency: A Laboratory Study on Communicating Risks in the Presence of Non-Events Stefan Traub (Helmut-Schmidt-Universität Hamburg)

Volunteering and Charity

TITEL: Recruiting Episodic Volunteers: A Field Experiment on Gift-Giving vs. Requesting a Favor

AUTOR*INNEN: Maximilian Hiller, Devin Kwasniok, Vanessa Mertins

ABSTRACT: Volunteers are indispensable for the effective functioning of nonprofit organizations (NPOs). Nevertheless, recruitment and motivation pose significant challenges. This study investigates the effectiveness of gift-giving and favor-requesting in enhancing volunteer engagement. A field experiment was conducted involving three groups: one group received a 10 Euro Amazon voucher for personal use, another group was asked to pass the voucher to an NPO, and a control group received neither. In both the gift-giving and favor-requesting groups, participants also had the alternative option to either keep the voucher for themselves or donate it to the NPO. The findings reveal no statistically significant differences between the treatments. This outcome carries important implications for NPO management, as implementing such treatments incurs costs. Furthermore, these results contrast with the typically positive evidence for reciprocity observed in laboratory studies.

TITEL: Boosting Performance Quality in Episodic Micro-Volunteering: A Collaborative Field Experiment

AUTOR*INNEN: Bernd Josef Leisen, Svenja Lange, Vanessa Mertins

ABSTRACT: The Impact of Monetary and Non-Monetary Incentives on Volunteer Performance Quality: Evidence from a Field Experiment on Micro-Volunteering for Elderly Care Facilities. Volunteers play a crucial role in various societal sectors, offering invaluable contributions that support both non-profit organizations (NPOs) and the broader social welfare system. Even micro time donations can have a significant impact if they are of sufficient quality, helping to alleviate the burden on professional staff and public social services. The challenge is how to ensure sufficient quality in the contributions of untrained one-shot volunteers. This paper contributes to the existing literature by examining the effects of monetary and non-monetary incentives on the performance quality of episodic volunteers through a field experiment. The study focuses on citizens who participated in a 5-minute, creative task: the development of digital recreational activities for elderly care facilities. The volunteers are randomly assigned to one of three groups prior to task execution: a control group, a group receiving a monetary incentive, and a group receiving public recognition of their contribution by name. The quality of their contributions is subsequently assessed by both the residents and professional caregiving staff of the elderly care facilities. The findings from this experiment provide insights into the relative effectiveness of different types of incentives in enhancing the quality of volunteer performance. This research has important implications for NPOs and policymakers aiming to optimize volunteer engagement and improve the support provided to the elderly population through volunteer initiatives.

TITEL: Fair Play: Leveraging Incentives and Games to Maximize Donations at a Volunteer Fair **CANCELLED!**

AUTOR*INNEN: Maximilian Hiller, Bernd Josef Leisen, Vanessa Mertins

ABSTRACT: Securing time, monetary, and in-kind donations is crucial yet challenging for nonprofit organizations, often impacting their ability to fulfill their missions. This study utilizes a unique setting to

conduct natural field experiments to enhance these essential resources within the nonprofit sector. Our research is conducted at the only volunteer-operated marquee at a large North German fair, which draws close to one million visitors over six days. We focus on approximately 100 unpaid episodic volunteers working in the tent and analyze how monetary and non-monetary incentives affect their willingness to commit to shifts. Revenues are primarily raised through a cake sale, and we examine the impact of these incentives on both the quantity and quality of cake donations. Furthermore, we explore the effectiveness of a specially designed donation box that makes a noticeable sound when money is deposited, investigating whether this public acknowledgment feature can increase monetary donations. Our findings aim to provide actionable insights into the mechanisms that can enhance volunteer participation and resource acquisition in the nonprofit sector.

AI and Humans II

TITEL: Can the Use of Large Language Models in Higher Education Empower Students to Increase their Learning Success?

AUTOR*INNEN: Eduard Buzila, Kai Heinrich

ABSTRACT: Large Language Models (LLMs) like ChatGPT have been reported to play an increasing role in higher education, particularly within the student's learning process. While it has been observed that LLMs can act as students passing university exams, there is no empirical evidence of the effects of LLM usage on students' learning success. Against this backdrop and utilizing the theory of Technology-Enhanced Learning and Bloom's updated taxonomy of educational objectives, we propose an experimental design to produce evidence of how well an LLM can act as a teacher in the context of higher education for different types of intended learning outcomes. Our experimental design is suppose to answer the question whether LLM can empower students to improve their learning success.

TITEL: Appreciation to biased algorithms, and aversion from unbiased algorithms.

AUTOR*INNEN: Ori Plonsky, Uri Hertz, Yefim Roth

ABSTRACT: Past research and common observations reveal conflicting evidence concerning people's tendency to appreciate or to reject advice from algorithms, particularly when compared to advice from human experts. Here, using three pre-registered experiments, we identify one contributing factor to the mixed evidence, namely that people learn to follow advice from a source that is accurate more often. Experiment 1 shows that humans who develop experience-based expertise often provide biased advice, recommending a course-of-action which is better most-of-the-time, but not on average. Experiments 2 and 3 show that advisees who receive this biased advice from human experts, as well as unbiased advice from algorithms, learn to follow the biased human advice – that is they learn to exhibit unbiased algorithm aversion. This happens because the unbiased algorithms give advice that, compared to the human expert

advice, is better on average but less frequently. This phenomenon replicates also when the advisors are explicitly labelled as humans or algorithms. Yet, when the algorithm is explicitly designed to be biased, and give advice that is better most-of-the-time, people learn to follow the algorithm more than they follow the human expert – that is they learn to exhibit biased algorithm appreciation. Our findings suggest humans prefer advice from sources that accommodate their own biases, and have both practical and ethical implications for the design of algorithms.

TITEL: Overreliance, Backfire, and the Importance of (Mis)trust in Human-Automation Interactions

AUTOR*INNEN: Doron Cohen, Yefim Roth, Markus Schöbel

ABSTRACT: As shared operational control between humans and automated systems becomes more common—such as in semi-autonomous cars, financial trading algorithms, and autonomous weapon systems—predicting and managing operator behavior is crucial. This paper investigates the conditions under which human operators over-rely on automated systems, leading to suboptimal risk-taking. We introduce an incentivized ‘Checking Dilemma’ where operators repeatedly decide whether to perform a costly redundant check or skip it, risking significant accidents. Our task design minimizes cognitive resource demand, focusing instead on incentives and feedback. We evaluate two automation system designs: Serial-checking, where the automated system’s decision is revealed to the human operator before their action, and Parallel-checking, where the decision is disclosed after the operator’s action. In Study 1 (N = 309), Parallel-checking led to better outcomes, while Serial-checking significantly increased overreliance and accident rates compared to a single-operator condition. In Study 2 (N = 205), when given a choice, operators preferred Serial-checking, resulting in significantly more overreliance. These findings suggest two main drivers: operators seek information that simplifies decision-making and tend to repeat actions that were mostly previously successful, even if detrimental in the long run. Understanding these behavioral drivers is essential for the effective design of automated collaborative systems.

Bargaining and Auctions

TITEL: The Economic Effects of Remote-Bargaining vs. Face-to-Face Bargaining

AUTOR*INNEN: Faeze Heydari, Claudia Keser, Holger A. Rau, Anne Schacht

ABSTRACT: In our study, we aim to examine the economic impact of remote negotiations (via Zoom) versus face-to-face negotiations through controlled laboratory experiments. Participants engage in a repeated variant of the Ultimatum Game with communication, negotiating the division of a sum of money. The stake of the negotiation is either low or high and known only to the proposer. Additionally, the study will employ facial-expression analyses to measure the emotions displayed by the negotiators during the two institutional settings. The results show that bargaining fails less often in remote negotiations, which leads

to higher welfare compared to the in-person treatment. While bargaining outcomes tend to be equitable in the majority of cases in the face-to-face negotiations, we observe higher inequity in the remote treatment.

TITEL: Between egalitarianism and utilitarianism: experimental evidence in unstructured bargaining

AUTOR*INNEN: Claus-Jochen Haake, Behnud Mir Djawadi

ABSTRACT: The literature on bargaining theory offers a wide range of solutions to the bargaining problem. We consider two non-symmetric two-person bargaining problems, each with constant local transfer rates (so-called hyperplane problems), and ask, whether participants in an unstructured bargaining experiment aim for an agreement that emphasizes individual or collective outcomes. While the former is described by the egalitarian bargaining solution, the latter corresponds to the utilitarian solution. We vary the transfer rates as well as the degree of anonymity of the bargaining partner and the possibility to execute side contracts after the experiment. In the experimental setting with lowest degree of anonymity and possibility of side contracts we predict that the utilitarian solution will be observed more frequently compared to all other settings. We further predict that the egalitarian solution will be observed more often the higher the degree of anonymity and the harder it is to side contract. We conduct the experiment both with students in a laboratory environment and with the broad population in a field setting within public events and city festivals. We have already collected more than half of the experimental data by now and will complete the data process for the student sample before the start of the conference. Further, the data collection of the broad population sample will hopefully be finished by then. Keywords: Bilateral Bargaining, Egalitarian Solution, Utilitarian Solution, Experiments JEL: C78, C91, C93, D63, D71.

TITEL: Lie Against AI: Revealing Private Information through AI in an Economic Experiment

AUTOR*INNEN: Nina Ostermaier, Dmitri Bershidsky, Laslo Dinges, Marc-André Fiedler, Jannik Greif, Ayoub Al-Hamadi, Joachim Weimann

ABSTRACT: Replicating the experiment of Belot & van de Ven (2017), we examine lying behavior in the presence of asymmetric information in a buyer-seller game. In our design, sellers have monetary incentives to sometimes misreport their private information about the color of a card assigned to them. Adapting the original study, we investigate the ability of buyers to detect lies via video conference and use the obtained video-communication to train a lie detection algorithm. Results indicate that sellers lie and buyers have limited ability to detect such lies. Further, we investigate the willingness of buyers to invest in revealing the private information of sellers using different methods, including the self-developed lie detection algorithm.

Social Norms and Reference Points

TITEL: Sanctions and Social Norms

AUTOR*INNEN: Sarah Necker, Matthias Kasper, **Amanda März**

ABSTRACT: This paper uses a novel empirical strategy to identify the causal effect of rule enforcement on social norms. Institutional rules can affect behaviour beyond providing material incentives by conveying their social meaning – as shown by the largely theoretical literature discussing the expressive power of the law. This literature argues that rules reflect, as well as shape, the social norms of a given society. However, the effect of rule enforcement on social norms remains largely unclear. In particular, while prior work establishes that the definition of rules affects social norms, the effect of discretion in the enforcement of rules on social norms remain unclear. Moreover, as societal rules, their enforcement, and norms typically co-evolve, establishing causality is difficult. We use a novel experimental design with a representative sample from Germany to identify the causal effect of rule enforcement on social norms. We aim to show that the enforcement of rules, rather than their definition, exerts expressive power and thus affects social norms.

TITEL: Social Networks and Coordination: A Simulation Experiment

AUTOR*INNEN: **Blanca Tena Estrada**, Swee-Hoon Chuah, Robert Hoffman, Chenting Jiang

ABSTRACT: In his 1973 article, Granovetter theoretically exposes how weak social link connections are the main driver of effective labor allocations. Information that is otherwise further away from one's social network, reaches further nodes through weak links. Thus, weak links facilitate coordination. Since then, multiple studies have empirically approached it, revealing evidence of their strength. Within our experiment, we aim to disentangle the different properties of strong vs weak link ties on coordination. Through a novel computer-simulated social structure of separate or overlapping groups, we study how the two types of links influence coordination in an online lab experiment. Participants play the minimum effort game with the same 4 other people 5 times per round, the strong ties network, or with 20 other people in five groups of 5 with no overlapping. We also add the treatment factor of communication via a restricted chat with the possible MEG integers

TITEL: Determinants and Dynamics of Collaborative Exploitation

AUTOR*INNEN: Thomas Meissner, Hannes Rusch, Julia Teufel, **Till Vater**

ABSTRACT: We designed an experimental framework that can be established as a benchmark paradigm for future research into labor exploitation on the level of individual principals and agents. The experiment consists of an extended gift exchange game and captures the inherent lack of contractibility in exploitative labor relationships. In parallel, we developed a survey module to capture an individual's propensity for both exploitative behavior and being exploited, mapping directly to the incentivized decisions in our

experiment. The experiment is conducted online via Prolific. The scope of the project is threefold. Firstly, we compare the observed behavior to the theoretical predictions. Secondly, we analyze the comparative statics of our experiment by changing the parameters of the game. Further, we investigate potential links between the observed exploitative behavior and personality traits, attitudes, and demographics. Our results show that the majority of participants behave exploitatively when given the opportunity. However, we observe gender differences in the severity of this behavior. Personal preferences correlate significantly with aspects of observed exploitative behavior

Health

TITEL: Human Infection Risk Assessment in Social Networks: The Role of Network Characteristics

AUTOR*INNEN: Martijn Stroom, Roselinde Kessels, Ingrid Rohde, **Martin Strobel**

ABSTRACT: Understanding an individual's human network assessment is fundamental for successful policy design when dealing with infection risk. Yet no research has successfully explored how humans cognitively mitigate their inability to handle complex network models. Using an online best-worst ranking experiment, we investigate individuals' perception of the risk of connecting to carefully constructed COVID-19-infected networks of 697 Dutch participants. We show that the perceived infection risk in social networks is not solely based on the objective probability of this risk: easily assessable physical characteristics have stronger predictive power on the perceived risk than the objective probability. Heterogeneity assessment suggests that demographic differences such as the level of education interact with the order and strength of the network characteristics. The often-complex mental calculation underlying objective risk in networks is substituted by a heuristics-driven approach. Our findings facilitate more tailored, effective, and generalizable policy and economically optimal infection-mitigating campaign design.

TITEL: An experiment on decision support in medicine

AUTOR*INNEN: Bernhard Roth, **Anna K. Stirner**, Daniel Wiesen

ABSTRACT: Purpose: Providing physicians with the option to use decision support to gather additional information for therapy recommendations can have a positive impact on quality of care, patient outcomes, and costs by ensuring that patients receive the care they need. However, a prerequisite for exploiting these positive effects is that healthcare providers make use of decision support options to make more appropriate decisions. We study how capacity constraints affect physicians' willingness to gather additional information supporting their therapy decisions. Further, we examine how capacity constraints affect the utilization of additionally gathered information and the appropriateness of physicians' therapy decisions. Methods: Using a controlled framed field experiment with German pediatricians (n=247), we exogenously vary the extent to which physicians' capacity is constrained. In our experiment, pediatricians make decisions on the length of antibiotic therapies for 40 pediatric routine cases. We use a between-subject design to vary two treatment parameters in our experiment: Availability of decision support and

information gathering costs. For each case, subjects first make an initial decision on the length of antibiotic therapy. Depending on the experimental condition they are assigned to, subjects then either (i) decide whether they want to use decision support before making the final therapy decision for that case, (ii) automatically get decision support, or (iii) do not have the option to use decision support. We vary the level of information gathering costs, which reflect the fraction of available capacities that is needed to use decision support. If a subject decides to use decision support or automatically receives support, he or she is given the opportunity to adjust his or her initial therapy decision for this case. Results: Our behavioral results evidence that physicians' willingness to gather additional information that supports decision making decreases as capacity constraints increase. However, the utilization of the information gathered is not affected by increasing constraints. We also find that capacity constraints have a statistically significant and clinically relevant impact on the appropriateness of therapy decisions and thus on the quality of care. This is especially the case for physicians with little clinical experience. The probability of gathering additional information increases significantly with case severity. Additional information increases the appropriateness of therapy decisions in the case of undertreatment, while overtreatment persists even after the use of decision support. Conclusions: Our behavioral results suggest that decreasing the extent to which capacity is constrained can be an effective way to enhance the utilization of decision support and thus help improve the appropriateness of therapy decisions. Implications of our findings for the management of healthcare organizations are discussed.

TITEL: Sow the Doubt, Reap the Inefficiency: A Laboratory Study on Communicating Risks in the Presence of Non-Events

AUTOR*INNEN: Jörg Felfe, Pauline Halm, Thomas Jacobsen, Jan Philipp Krügel, Yvonne Nestoriuc, Fabian Paetzel, Mey-Ling Sommer, **Stefan Traub**

ABSTRACT: An important task of public health agencies is to inform the public about health risks. Preventive measures and changes in people's behavior can reduce both the number of people affected and the individual damage to health, so that the actual impact remains lower than predicted. The literature suggests that people who have experienced such a non-event ask themselves in retrospect whether the health risk was really as high as communicated and whether the preventive measures and behavioral changes, which are usually associated with costs, were sensible. We provide a theoretical framework for the behavioral economic analysis of individual risk-taking in the presence of non-events and test it in a laboratory experiment. Risk communication generally increases the efficiency of risk-taking. Communication of aleatory or epistemic uncertainty has no adverse effect on individual risk taking while calling into question the credibility of the health agency significantly deteriorates the efficiency of risk communication.

Session 5 (Freitag 27.09., 09:00 Uhr)

S11	S12	S14	S15
Ethics and Fairness II (Sessionsprache: Deutsch)	AI, Transparency and Bias	Cooperation and Prosociality III	Field Experiments
Markets, Social Responsibility and the Replacement Logic Janina Kraus (TU Clausthal)	Literature review on Explainability of Artificial Intelligence from a Behavioral Perspective (Working title) Anika Bittner (Technische Universität Hamburg)	Image Concerns and Reputation Fabian Hoffmann (Universität zu Köln)	Experimenting with Financial Professionals Christian König genannt Kersting (Universität Innsbruck)
The Impact of Tax Culture on Tax Rate Structure Preferences: Results from a Vignette Study with Migrants and Nonmigrants in Germany Dirk Kieseewetter (Julius-Maximilians-Universität Würzburg)	How does being involved in AI training affect expert user adherence - An experiment in the context of AutoML Marius Protte (Universität Paderborn)	Institutional Change in the Infinitely Repeated Prisoners' Dilemma Christine Meemann (Helmut-Schmidt-Universität Hamburg)	Do AI skills improve employment chances? Evidence from a field experiment# Anastasia Danilov (Humboldt-Universität zu Berlin)
Fairness Deliberations and Fair Allocations in Symmetric and Asymmetric Bargaining An Experimental Study on Group Decisions in Germany and China Heike Hennig-Schmidt (University of Bonn)	Measuring social bias in face perception of a vision-language model Carina Hausladen (ETH Zurich)	Inflated Rules in Threshold Public Goods Games Philipp Monschau (RWTH Aachen)	Motivating the Weak: The Impact of Group-Based Incentives on Teams Devin Kwasniok (Universität Vechta)

Ethics and Fairness II

TITEL: Markets, Social Responsibility and the Replacement Logic

AUTOR*INNEN: Miguel Abellán, Janina Kraus und Mario Mechtel

ABSTRACT: In accordance with the rationale "if I don't, somebody else will," participants in the market may opt to engage in a trade that increases their monetary payoff rather than allowing another party to do the same (e.g., Sobel 2010). Building on the experimental market paradigm in Bartling et al. (2015, 2019) and on previous research by Bartling and Özdemir (2023) and Ziegler et al. (2024), the present study investigates whether buyers engage in a trade that increases their monetary payoff but harms a third party under the rationale that otherwise another buyer may engage in the same trade and cause the harm anyway. In order to test the hypothesis that the replacement logic affects social responsibility in markets negatively, we compare three treatments: a treatment condition with one buyer (T1, baseline), two buyers (T2, low competition) and four buyers (T4, high competition) in the market. In T2 and T4, buyers enter the market sequentially only if the buyer(s) who has (have) previously entered the market did not buy the product. Our findings provide evidence for the argument that the replacement logic affects social responsibility in markets negatively. First, the overall trade frequency of the product with a negative externality on the third party is higher the less pivotal buyers are (93.7% in T4 85.2% in T2 57.8% in T1). Second, the trading frequency of this product by the first buyer in the market is also higher the less pivotal buyers are (73.4% in T4 70.3% in T2 57.8% in T1)."

TITEL: The Impact of Tax Culture on Tax Rate Structure Preferences: Results from a Vignette Study with Migrants and Nonmigrants in Germany

AUTOR*INNEN: Dirk Kieseewetter, Andre Machwart

ABSTRACT: We explore the relationship between tax culture and tax rate structure preferences among migrants and nonmigrants in Germany. A vignette study is used to examine (1) whether migrants bring their country of origin's tax culture to the destination country and (2) whether second-generation migrants assimilate with the host society's tax culture. Our findings provide evidence for the impact of tax culture. Migrants tend to prefer a less-progressive tax rate structure, especially those from flat tax countries. Additionally, while second-generation migrants align their preferences with those of the host society, differences remain. This research provides insights into the dynamics of tax culture in heterogeneous societies.

TITEL: Fairness Deliberations and Fair Allocations in Symmetric and Asymmetric Bargaining An Experimental Study on Group Decisions in Germany and China

AUTOR*INNEN: Heike Hennig-Schmidt, Zhuyu Li, Gari Walkowitz

ABSTRACT: The study's primary focus is on examining the interaction between fairness concerns and power asymmetry to understand whether bargainers in Germany and in China incorporate fairness into their decision process and, if so, how this affects bargaining outcomes. To this end, we conducted an incentivized ultimatum bargaining experiment with symmetric and asymmetric outside-options. Groups (N=142) of three persons interact as proposers and responders in dyads and decide simultaneously on their offer or which offers to accept or reject, respectively. Communication between parties is inhibited. We videotaped in-group discussions the resulting transcripts were text analyzed by eliciting whether groups make fairness an issue, which fairness norms they discuss, and whether they use fairness-related perspective-taking to overcome the communication constraint. We find that asymmetry of bargaining power in favor of the proposer induces lower offers relative to the symmetric situation. Fairness concerns alone have no significant impact on offers. However, when associated with the equal-payoff norm, and in Chinese groups in particular, discussing fairness increases offers in symmetric but also in asymmetric conditions, in which other fairness norms could have been applied, too. Fairness-related perspective-taking is used by German and Chinese groups and is associated with higher offers in the former. Our study makes an epistemological and related methodological contribution: a possibly biased interpretation of bargaining outcomes can be avoided if information on decision processes and underlying mechanisms were available.

AI, Transparency and Bias

TITEL: Literature review on Explainability of Artificial Intelligence from a Behavioral Perspective

AUTOR*INNEN: Anika Bittner, Timo Heinrich

ABSTRACT: TBA

TITEL: How does being involved in AI training affect expert user adherence - An experiment in the context of AutoML

AUTOR*INNEN: Anastasia Lebedeva, Marius Protte, Dirk van Straaten, René Fahr

ABSTRACT: AutoML is a promising field of Machine Learning (ML) that is supposed to bring the advantages of artificial intelligence to a wide range of organizations in plentiful domains, by automating the process of ML-model creation without requiring prior knowledge in data science or programming. However, AutoML often appears to users as a black-box model created in a black-box process, negatively impacting users' trust. Additionally, AutoML users are often experts in their respective domains (physicians, engineers, etc.), which are commonly observed to exhibit stronger algorithm aversion than lay people, i.e., having more difficulties trusting and relying on AI recommendations. User non-adherence to AutoML may have high-

cost consequences, resulting in inefficient decisions and mitigating the overall progress in the AutoML field. Therefore, we investigate how domain experts' adherence to AutoML recommendations can be fostered. As involvement of users in product creation processes was shown to positively affect their attitudes towards the product in multiple contexts, we argue that involving domain experts in AutoML-model creation processes may increase their trust and adherence to AutoML. We conduct an experimental laboratory study, in which subjects act as expert engineers and need to foresee machine malfunctions, while being advised by an AutoML-model. We apply three treatments – zero, passive & active involvement – to investigate our hypothesis. We observe that higher involvement leads to a higher perceived influence on the AutoML model and a higher perceived understanding of its functionality. However, these perceptions are not reflected in the actual behavior – subjects across all groups demonstrate similar AI adherence.

TITEL: Measuring social bias in face perception of a vision-language model

AUTOR*INNEN: Carina I. Hausladen, Manuel Knott, Colin F. Camerer, and Pietro Perona

ABSTRACT: Measuring algorithmic bias is the first step towards deploying AI systems responsibly. We introduce a new method to assess biases in the perception of human faces by vision-language models. Our method represents a novel approach by integrating three key elements: direct analysis of embeddings, the application of well-calibrated terms from social psychology, and the use of an experimental, synthetic face dataset. We apply our method to measure the social biases of face perception in CLIP, the most popular open-source embedding. We quantify social judgments by computing the embedding similarity between face images and valence words from the two leading social psychology theories of human stereotypes. The face images are varied systematically and independently with respect to legally protected attributes such as age, gender, and race, as well as a number of additional visual attributes, such as facial expression, lighting, and pose. Thus, our tests are experimental rather than observational and allow us to study the effect of varying specific attributes one at a time and avoid spurious correlations. Our experiments reveal that variations in the protected attributes systematically impact the model's social perception of faces. Additionally, we find that facial expression impacts the model's social perception as much as gender and more so than age. Thus, our findings highlight that controlling for non-protected attributes is necessary when assessing bias vis-a-vis protected attributes. Moreover, we find that faces within each demographic group are rated similarly in warmth and competence, while these perceptions significantly differ when comparing across different demographic groups. Most strikingly, our approach highlights patterns of bias that are peculiar to the faces of Black women, where CLIP produces extreme values of social perception across different ages and facial expressions. We compare our findings with those produced through the traditional observational method using two public datasets of face images. In this case, we find that the distinction in perceptions of different demographic groups is no longer significant, suggesting that measurements using observational wild-collected datasets are more noisy due to uncontrolled variations of attributes such as lighting, facial expression, viewpoint, and background.

Cooperation and Prosociality III

TITEL: Image Concerns and Reputation

AUTOR*INNEN: Fabian Hoffmann

ABSTRACT: Public recognition typically increases prosocial behavior because of a desire to be well regarded by others. I experimentally investigate how reputations of low or high prosociality affect the responsiveness to public recognition. At the aggregate level, a reputation of high prosociality significantly decreases image-relevant prosocial behavior compared to a situation with no reputation. A reputation of low prosociality tends to increase image-relevant prosocial behavior, but not significantly. The effect sizes are small due to a substantial behavioral heterogeneity between participants. Opposing behavioral responses largely cancel each other out. When a reputation of high prosociality is public, 74 percent of participants with a non-zero treatment difference decrease their prosocial behavior, while 26 percent increase their prosocial behavior. When a reputation of low prosociality is public, 42 percent of participants with a non-zero treatment difference decrease their prosocial behavior, while 58 percent increase their prosocial behavior. Participants' treatment responses can be explained by their social image concerns and intrinsic prosociality.

TITEL: Institutional Change in the Infinitely Repeated Prisoners' Dilemma

AUTOR*INNEN: Lydia Mechtenberg, Christine Meemann, Stefan Traub

ABSTRACT: Payoffs that follow from collective action often depend on this collective action themselves, e.g., collective behavior induces climate change that, in its turn, affects the payoffs resulting from it. We experimentally study subjects' cooperation behavior in the infinitely repeated prisoners' dilemma with reinforcement. Using the model setup by Greif and Laitin (2004), we test whether positive feedback (cooperation increases future payoffs) leads to higher cooperation rates and negative feedback (cooperation decreases future payoffs) leads to lower cooperation rates as compared to neutral feedback. Furthermore, we differentiate between two cases: In the first round of the infinitely repeated prisoners' dilemma, mutual cooperation is either a sub-game perfect equilibrium or not.

TITEL: : Inflated Rules in Threshold Public Goods Games

AUTOR*INNEN: Philipp Monschau, Christian Grund

ABSTRACT: Most people would approve that living together demands behavioral rules. Those rules shall steer actions of individuals to secure an outcome that is best for the society. Often, such rules are "inflated" in the sense that they are requesting a behavior which is excessively demanding the targeted goal could be reached also by a more lenient rule i.e. a manager asks her employees to double or triple check their own work or that of co-workers. Considering the average behavior, this "inflation" can be seen as a validation against accidental inattention or conscious rule violation of individuals. But this additional demand in behavior also opens up a space for beneficial deviations from following the rule. Our experimental design

shall give insights into the structure of beneficial behavioral rules. Participants will face a threshold public goods game. The public goods situation models a typical cooperation dilemma where group interests collide with private preferences. The integration of a threshold adds a coordinating aspect to the task and through this, allows a closer look on efficiency in terms of including (i) the likelihood to reach the threshold and (ii) the degree of excessive contributions. By limiting the feasible contributions to integers and asking for a threshold that cannot be split equally among the participants, an optimal partition of the contributions is not obvious to the players. A behavior rule that prescribes a partition could work as norm induction as well as coordination device. Our design contains different treatments which vary (i) in the amount needed of four subjects for the publicly announced joint threshold and (ii) the announced rule in terms of strong recommendations towards individuals to stick to a demanded contribution to the public good. Here, we distinguish situations, in which rules do or do not differ across individuals in groups and rules which sum up exactly to the threshold or are inflated in the sense that in sum contributions beyond the threshold are demanded. We want to address the following research questions: 1.) Do inflated rules lead to more frequently reached thresholds? 2.) Do inflated rules lead to a waste of resources in terms of exceeding contributions? 3.) What effect do unequal rules across individuals in groups have for questions 1 and 2.

Field Experiments

TITEL: Experimenting with Financial Professionals

AUTOR*INNEN: Christoph Huber, **Christian König-Kersting**, Matteo M. Marini

ABSTRACT: As key players in financial markets and the broader industry, financial professionals are increasingly used as experimental research participants. We review over 50 studies comparing financial professionals to laypeople and conduct systematic meta-analyses of 24 eligible studies spanning from 1986 to 2023. Our findings reveal persistent and robust support for financial professionals being more risk- and uncertainty-loving, but little evidence of superior forecasting accuracy. Further analyses indicate that larger monetary payments result in greater behavioral differences between financial professionals and laypeople, suggesting an increased susceptibility to incentives among professionals. This systematic review not only synthesizes experimental results, contributing to recent discussions about external validity and generalizability, but also highlights critical methodological considerations when experimenting with financial professionals.

TITEL: Do AI skills improve employment chances? Evidence from a field experiment

AUTOR*INNEN: **Anastasia Danilov**, Teo Firpo, Lukas Niemann

ABSTRACT: With the advent of Artificial Intelligence (AI), a growing number of firms are adopting this technology, leading to a rise in demand for AI-related skills. There is, however, little evidence on the impact of acquiring AI skills for prospective job seekers. We address this question with a correspondence study. We send résumés, randomizing whether they contain AI skills, to 1,185 entry-level job ads in the UK across

a range of job functions. Overall, we find no evidence that including AI skills in a résumé increases a candidate's chances of being invited for an interview. Nevertheless, we observe some heterogeneity by job function, where the inclusion of AI skills in résumés has a positive effect on call-back rates for Engineering and Marketing jobs. We explore alternative explanations with the textual analysis of the vacancies and a survey of professionals with hiring experience.

TITEL: A No title provided! The paper is about a field experiment on team

AUTOR*INNEN: Maximilian Hiller, Devin Kwasniok, Vanessa Mertins

ABSTRACT: We examine the effects of different group-based incentive schemes on team performance regarding step counts in a framed field experiment. Specifically, we explore the impact of minimum (maximum) incentives, where the weaker (stronger) partner of the day becomes wholly or in half payoff relevant for the team. While heterogeneous teams show significant performance gains across most treatments, the benefits are particularly pronounced when the payoff solely depends on one team member. Additionally, most treatments significantly enhance the team performance of participants who initially walked less than the daily maximum step goal. This suggests that tailored incentive schemes are effective for improving performance among low-performers. However, the performance of stronger participants, especially those who walked more than the maximum goal before the intervention, decreases when partnered with weaker teammates. Goal setting and individual capabilities must be carefully balanced to maximize overall team performance.

Session 6 (Freitag 27.09., 11:00 Uhr)

S11	S12	S14
Communication	Freedom and Authority	Teamwork
Experimentally examining multiple normative expectations in whistleblowing and the influence of social information interventions Sabrina Plaß (Universität Paderborn)	Giving a Voice - Increasing Individual Self-Expression to Enhance Group Welfare and the Resilience to System Disbelief Mathilde Draeger (Otto-von-Guericke-Universität Magdeburg)	Social Closeness in Coordination Games with Conflict Kevin Grubiak (Universität Passau)
Misrepresenting Risks: A sender-receiver model for efficient outcomes Mey-Ling Sommer (Helmut Schmidt Universität)	Knowledge and Freedom Max R. P. Grossmann (Universität zu Köln)	Creative Performance in Teams: The Effect of Workplace Settings and Self-Selection - an Experimental Study Lisa Klinkenberg (RWTH Aachen)
Avoid, recycle, compensate - An experimental study on nudging environmental behavior Marina Schröder (Leibniz Universität Hannover)	Authoritarian Preferences Gönül Doğan (Universität zu Köln)	It's a match! Team formation and performance in innovation-related tasks Nathalie Römer (Leibniz Universität Hannover)

Communication

TITEL: Experimentally examining multiple normative expectations in whistleblowing and the influence of social information interventions

AUTOR*INNEN: Behnud Mir Djawadi, **Sabrina Plaß**, Sabrina Schäfers

ABSTRACT: Internal whistleblowing is often considered a trade-off between establishing fairness by stopping wrongdoing and reporting a colleague to whom one feels loyalty. Therefore, the whistleblowing decision may be influenced by multiple normative expectations that cannot be met simultaneously because a potential whistleblower might perceive that others find both behavioural options, whistleblowing and remaining silent, similarly (in-)appropriate. Hence, in our study, we measure normative expectations regarding both behavioural options (whistleblowing and remaining silent) and investigate how they (jointly) affect the whistleblowing decision. Based thereon, we design social information interventions to examine if/how communicating normative expectations about the appropriateness of both behaviours (whistleblowing and remaining silent) compared to communicating just one increases whistleblowing behaviour. We investigate this by means of an incentivized experiment (basic experiment and three treatments) that we conduct on the platform Prolific with real employees.

TITEL: Misrepresenting Risks: A sender-receiver model for efficient outcomes

AUTOR*INNEN: Mey-Ling Sommer

ABSTRACT: This paper investigates efficiency outcomes in a two-period sender-receiver model. It has been theoretically shown that there is an optimal level of downplaying risks if receivers have adaptive expectations on lottery outcomes. The model is tested in a laboratory online-experiment with students from the University of Hamburg. Participants are randomly assigned into groups of five and decide individually whether to play a lottery. Deviating beliefs of receivers regarding the expected lottery outcomes in the first period may trigger a more risky behavior in the second period leading to lower efficiency. It is hypothesized that downplaying the communicated risk level in the lottery may counteract this efficiency loss.

TITEL: Avoid, recycle, compensate - An experimental study on nudging environmental behavior

AUTOR*INNEN: Annika Herr, Soschia Karimi, **Marina Schröder**

ABSTRACT: We analyze the effect of information nudges encouraging waste separation on waste avoidance, waste separation, and compensation for produced waste. In our experimental study, participants conduct a task that produces different types of waste. Across treatments, we vary whether or not a nudge to promote waste separation is provided to participants. Furthermore, we vary if the nudge emphasizes the social or the personal benefits of waste separation. We find that nudging can lead to a significant increase in waste separation, and that nudging does not lead to undesirable

spillover effects on other dimensions of sustainable behavior, namely wasteavoidance or waste separation.

Freedom and Authority

TITEL: Giving a Voice - Increasing Individual Self-Expression to Enhance Group Welfare and the Resilience to System Disbelief

AUTOR*INNEN: Mathilde Draeger

ABSTRACT: Individuals in a group, who repeatedly experience that their group's policy selection system does not decide in their favor, often develop system disbelief. The notion that the system is not favorable for the group, i.e., system disbelief, may be detrimental to the performance and welfare of the group. It may affect group members' psychological well-being and their willingness to provide work effort, make financial contributions, or participate in cooperative coordination. In this experimental study, I investigate whether giving individuals a voice, i.e., the opportunity to express and explain their preferences, mitigates the development of system disbelief when decisions are made by AI (ChatGPT) or by an independent human committee. The study adds to the knowledge on the drivers of discontent in group decision processes. It provides insights for managers and policymakers concerning the design of group decision processes that are welfare enhancing and resilient towards system disbelief.

TITEL: Knowledge and Freedom

AUTOR*INNEN: Max R. P. Grossmann

ABSTRACT: We study the relationship between information and paternalism. When is autonomy granted to an individual that is less or better informed, and if no autonomy is granted, what form will intervention take? We introduce the Estimation Game, a novel experimental design that varies the amount of ambiguity inherent in a binary lottery. Our analysis is concerned with the behavior of policymakers ("Choice Architects"). We conduct experiments of the Estimation Game with a Choice Architect to examine the extent to which Chooser knowledge matters in terms of restrictions on the Chooser's freedom of choice. More information leads to fewer interventions on the extensive margin. However, which option is imposed is a matter of personal Choice Architect preference, not a counterfactual assessment of the other's choice. A followup experiment replicates these results, highlighting that Choice Architects rely on their own preference when choosing the intensive margin. However, when Choice Architects are informed about the Chooser's hypothetical full-information choice, this information is used to determine what to impose. Moreover, Choice Architects disproportionately want the Chooser to decide informedly, even when Choice Architects are able to exploit the Chooser's ignorance.

TITEL: Authoritarian Preferences

AUTOR*INNEN: Gönül Doğan, Louis Strang

ABSTRACT: Right-wing populist movements favour authoritarian decision-making rules that would give little to no decision-making powers to minorities. One purported reason for increasing support for populist discourse is identity threat. In this project, we probe whether minority success generates such an identity threat to the majority group thereby increasing discrimination as well as support for authoritarian decision rules. In a survey experiment on a representative German sample varying exposure to minority success stories, we find substantial support for authoritarian decision-rules in a resource allocation task with very high stakes regardless of the treatment. Groups without no or few minorities are preferred over other groups while allocating decision power. Exposure to minority success stories have no effect in the representative sample. East and West Germany differ in the directional effect of the treatment in various measures: Whereas in East Germany, discrimination towards minorities and support for authoritarian rules increase with exposure to minority success stories, in the West, they pointwise decrease, resulting in a treatment effect in the difference between the two regions.

Teamwork

TITEL: Social Closeness in Coordination Games with Conflict

AUTOR*INNEN: Kevin Grubiak

ABSTRACT: The power of focal points is often found to be limited in games involving conflicts-of-interest. Previous studies have suggested that the frequent miscoordination in these games is due to a tendency to use individualistic level-k reasoning instead of collective team reasoning. We experimentally investigate whether socially close participants are more likely to overcome conflicts by thinking and acting as a team. To gain insights into participants' reasoning processes, we use an intra-team communication protocol.

TITEL: Creative Performance in Teams: The Effect of Workplace Settings and Self-Selection - an Experimental Study

AUTOR*INNEN: Lisa Klinkenberg, Christian Grund, Christine Harbring

ABSTRACT: The organization of work and the characteristics of tasks have considerably changed over recent years. The developments include (i) an increased relevance of virtual teams and (ii) a higher demand for creative work in organizations, also in teams. One crucial challenge re-garding effective

team coordination can be seen in the initial stage of a team. We argue that this is supposed to be relevant in the virtual case and for creative work in particular. Therefore, our research focuses on the creative performance in newly-formed teams and whether there is a difference between work settings, i.e. on-site and virtual teams, and if the sequence of settings is of importance. We also address the question of whether the individuals' ability to decide where to work matters for the creative performance of teams. We answer those questions by conducting a 2-phase experiment in the lab and online to model an on-site and a virtual work setting. To simulate a creative team task, we implemented the "Unusual-Uses Task" in a dyadic team setting. We implemented four treatments of dyadic teams, in which we also varied the work setting over two sessions between virtual and lab. Additionally, we implemented two treatments where participants (endogenously) decided on the work setting in advance. Our results showed that working at least one phase in presence leads to higher creative performance than solely working online. Moreover, no significant effects of self-selection on the creative performance are found. Therefore, when implementing virtual teams, the workplace setting during the first interaction needs to be carefully selected.

TITEL: It's a match! Team formation and performance in innovation-related tasks

AUTOR*INNEN: Joshua Graff Zivin, Nathalie Römer

ABSTRACT: We provide causal evidence of the effect of social connections on team formation preferences. Using a novel experimental design, we induce social connection in an online study via a short 2 minutes video conversation. We can show that workers prefer to form a team with a lower-skilled worker they have communicated with prior to the team task than with higher-skilled workers they have not spoken to. By examining team performance, we show that preference-based matching does not result in better outcomes. Our findings imply that self-formed teams may be configured sub-optimally as they can be biased by social connections.