

ACL読み会 @名大

幹事

名大・笹野研

D1 大鹿, D1 矢野

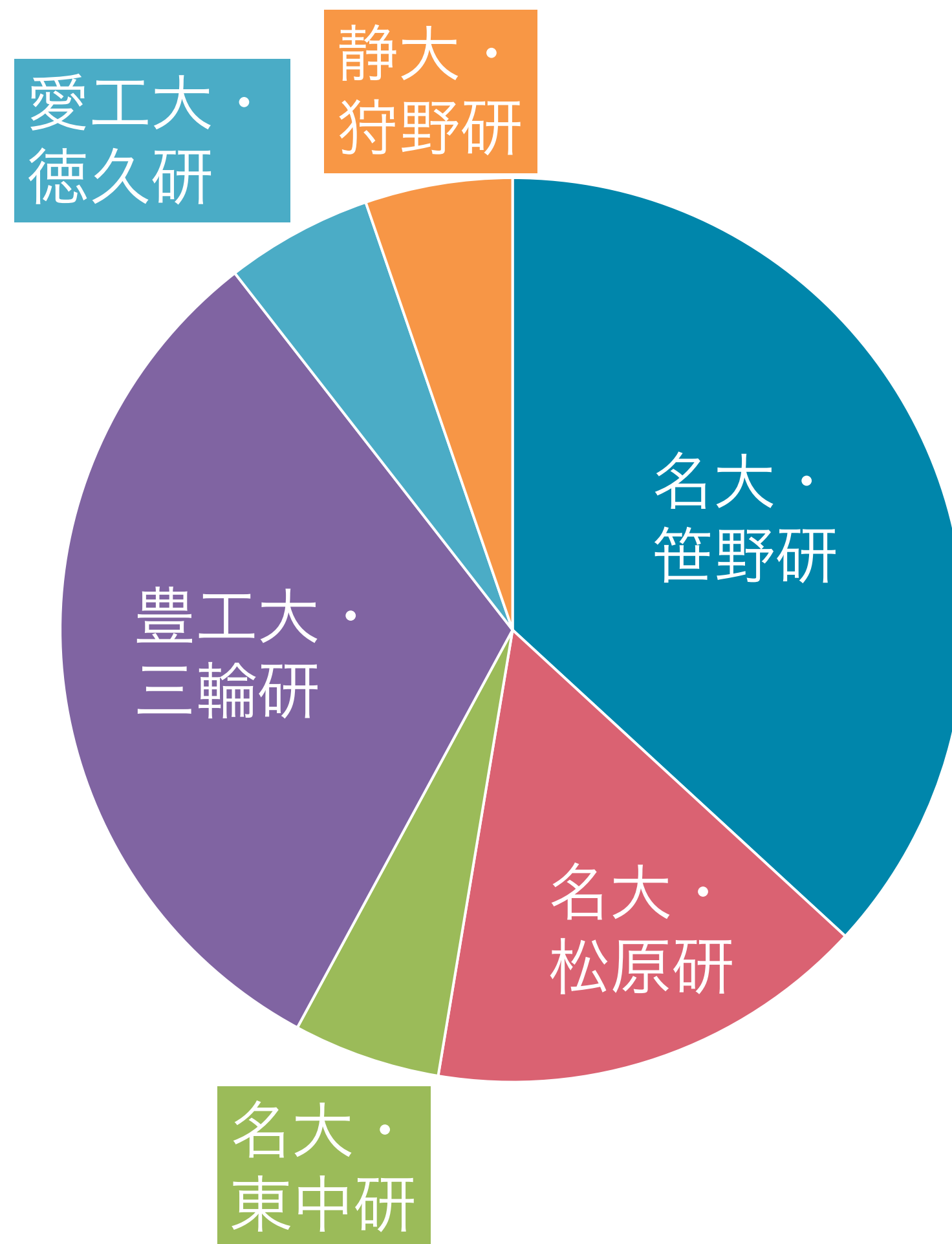
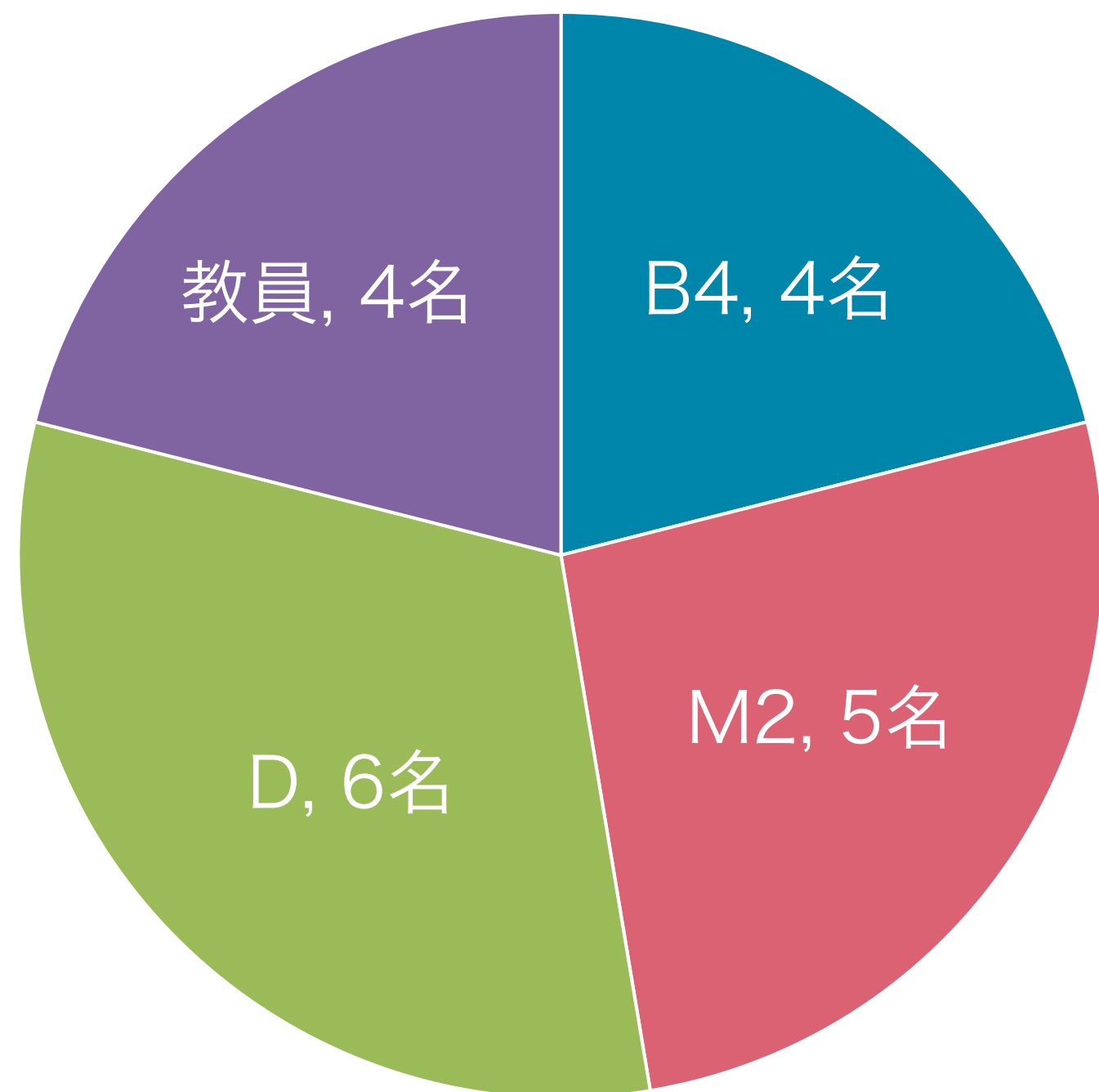
プログラム

- 2, 3人で1セッション
 - 1セッション36分 or 55分
- 1人18分
 - 発表15分質疑3分
- セッション間10分
- 昼休み55分
- 中休み約20分

プログラム

- 9:30-9:35: Opening [大鹿, 矢野 [slide](#)]
- 9:35-10:30: Session1: LLM:応用
 - [Evaluating Language Models as Synthetic Data Generators](#) [大鹿, [slide](#)]
 - [Completing A Systematic Review in Hours instead of Months with Interactive AI Agents](#) [徳久, [slide](#)]
 - [Faithful and Robust LLM-Driven Theorem Proving for NLI Explanations](#) [加藤, [slide](#)]
- 10:40-11:35: Session2: データセット、ベンチマーク
 - [ConLoan: A Contrastive Multilingual Dataset for Evaluating Loanwords](#) [田中, [slide](#)]
 - [RAGEval: Scenario Specific RAG Evaluation Dataset Generation Framework](#) [小島, [slide](#)]
 - [Can Knowledge Graphs Make Large Language Models More Trustworthy? An Empirical Study Over Open-ended Question Answering](#) [黄, [slide](#)]
- 11:35-12:35: Lunch Break
- 12:35-13:30: Session3: モデル理解
 - [The Impact of Token Granularity on the Predictive Power of Language Model Surprisal](#) [笹野, [slide](#)]
 - [A Theory of Response Sampling in LLMs: Part Descriptive and Part Prescriptive](#) [井田, [slide](#)]
 - [Understanding the Dark Side of LLMs' Intrinsic Self-Correction](#) [張, [slide](#)]
- 13:40-14:35: Session4: 埋め込み表現
 - [On the Acquisition of Shared Grammatical Representations in Bilingual Language Models](#) [山口, [slide](#)]
 - [Length-Induced Embedding Collapse in PLM-based Models](#) [矢野, [slide](#)]
 - [Mapping 1,000+ Language Models via the Log-Likelihood Vector](#) [木下, [slide](#)]
- 14:35-14:55: Short Break
- 14:55-15:31: Session5: 文書理解
 - [The Noisy Path from Source to Citation: Measuring How Scholars Engage with Past Research](#) [伊藤, [slide](#)]
 - [Hierarchical Document Refinement for Long-context Retrieval-augmented Generation](#) [韓, [slide](#)]
- 15:40-16:16: Session6: LLM:推論
 - [Large Language and Reasoning Models are Shallow Disjunctive Reasoners](#) [大崎, [slide](#)]
 - [Can Language Models Reason about Individualistic Human Values and Preferences?](#) [錢本, [slide](#)]
- 16:16-16:40: Short Break
- 16:40-17:35: Session7: モデル効率化
 - [500xCompressor: Generalized Prompt Compression for Large Language Models](#) [井上, [slide](#)]
 - [Byte Latent Transformer: Patches Scale Better Than Tokens](#) [佐藤, [slide](#)]
 - [Unifying Continuous and Discrete Text Diffusion with Non-simultaneous Diffusion Processes](#) [金児, [slide](#)]
- 17:35-17:40: Closing

発表者（19名）の内訳



進め方・お願い

- 質疑等はnu-nlpワークスペース内の#acl読み会2025にて行います
 - 参加されていない方は参加の方よろしくをお願いいたします
- 発表者の方は接続チェック＆スライドの共有をお願いします
 - 接続チェックは自分のセッションの前の休み時間をお願いします
 - スライドは可能であればWebページで共有させていただきたいです
 - ホームページにて資料の共有をしています
- 質問は口頭またはSlackをお願いします
 - 発表者ごとにスレッドを作成するのでご活用ください
 - 質問がない場合はどんどん進めていきます