

# AI Gender Bias in Moral Guidance: A Computational Content and Sentiment Analysis of ChatGPT, Gemini, Le Chat, and DeepSeek

*Thursday 18 September 2025 15:15 (15 minutes)*

We present a computational study of gender bias in four state-of-the-art large language models (LLMs), namely ChatGPT, Gemini, Le Chat, and DeepSeek, using a novel prompt-based framework to evaluate model responses to moral guidance questions. Our methodology integrates structured prompt variation, content analysis, sentiment scoring, subjectivity detection, part-of-speech tagging, and n-gram extraction. Original contributions include: (1) a cross-model comparison of LLM outputs to gendered and neutral prompts within a controlled experimental setup; (2) quantitative evidence of systematic lexical, affective, and structural variation based on prompt gendering; and (3) a reproducible hybrid method combining sociolinguistic analysis with automated metrics for bias detection. Results indicate that gendered prompts yield distinct affective profiles and lexical distributions, with female-coded inputs eliciting more nurturing and communal language and male-coded prompts emphasizing autonomy and discipline. All models exhibit consistently positive sentiment, but differ in subjectivity and stylistic framing. These findings show persistent gender-normative patterns in moral guidance outputs and demonstrate the utility of integrating qualitative reasoning with NLP techniques for auditing LLM behavior. We argue that prompt-sensitive evaluation is essential for assessing alignment and mitigating subtle forms of discursive bias in generative systems.

**Authors:** Ms IORDACHE, Diana-Alexandra (Faculty of Sociology and Social Work, University of Bucharest); Dr SIMINIUC, Rodica (Technical University of Moldova)

**Presenter:** Ms IORDACHE, Diana-Alexandra (Faculty of Sociology and Social Work, University of Bucharest)

**Session Classification:** Data FAIR in Science

**Track Classification:** Social Aspects of Networking Environment Today