

Here are 10 comprehensive bullet points from the seminar transcript, each with a one-sentence summary.

- **Introduction to AI via Nobel Laureates** The seminar began by highlighting the 2024 Nobel Prize winners Demis Hassabis and Jeffrey Hinton, whose work in cognitive science and backpropagation (a key AI training concept) is foundational to modern AI development .
 - *Summary:* The session's start connected today's AI advancements directly to Nobel Prize-winning work in cognitive science.
- **Seminar Goal: Moving to "Regular Use"** The instructor's stated goal is to move participants beyond simple "awareness" of AI (Level 1) to "regular use" (Levels 2 and 3), where they can confidently apply AI for practical tasks like document summarization, writing, and even coding .
 - *Summary:* This seminar aims to make attendees skilled, practical users of AI, not just people who know what it is.
- **The Technological Drivers of AI** Today's powerful AI is only possible due to exponential growth in three key areas: processor power (billions of transistors), data transmission speed (gigabit internet), and the massive scale of AI "parameters," which now number in the trillions .
 - *Summary:* AI's current power is built on massive, decades-long increases in computing power, internet speed, and data processing capability.
- **Core AI Vocabulary Explained** Several key terms were defined, including **Prompt** (the request you make), **Inference** (the AI's answer), **RAG** (Retrieval Augmented Generation, which lets AI get current web info), and **RLHF** (Reinforcement Learning from Human Feedback, the process of humans training the AI) .
 - *Summary:* The group learned a new vocabulary to understand how AI works, including how it gets answers and how it is trained by people.
- **AI and "Theory of Mind"** A presentation (made by an AI) explored the "Theory of Mind," discussing how AI is now interpreting human speech, tone, and pace to infer emotional states, while humans, in turn, are beginning to perceive AI as conscious .
 - *Summary:* The class examined the complex idea of AI "reading" human emotions and, conversely, humans forming emotional bonds with AI.
- **A 4-Part Structure for Effective Prompts** To get the best results, the instructor recommended a four-part prompt: assign a **Role** ("You are an expert..."), provide **Context** (the background), state the **Task** (the question), and define the **Format** ("in five bullet points") .

- *Summary:* A clear, four-step method (Role, Context, Task, Format) was taught to help everyone write better prompts.
- **Live Demonstrations of AI Tools** The instructor demonstrated multiple AIs live, using Adobe Firefly to generate images, Grok to create a complete business plan, and Meta AI to create images from a mathematical formula, showing the different strengths of each tool .
 - *Summary:* Practical, live examples showed how different AI tools can be used for creative, business, and even academic tasks.
- **The Challenge of "Markdown" Formatting** When an AI's response (like the business plan from Grok) was pasted into Google Docs, it included messy formatting code, which the instructor identified as **Markdown**—a common problem users will face .
 - *Summary:* A common technical problem was shown: copying text from AI can result in messy "Markdown" code that users need to learn to handle.
- **Inherent AI Biases and Homogeneity** AIs have built-in biases based on their training data (e.g., models trained on English may skew "liberal"), and if prompts are not specific, the AI will default to the most common, or "homogeneous," answer based on statistics .
 - *Summary:* The instructor warned that AI answers are shaped by their data, which can lead to biased or overly generic results.
- **Fact-Checking and AI Hallucinations** In the Q&A, it was explained that AIs can "hallucinate" or make up facts because they are "aligned to please you"; the best defense is to use your own human judgment or, as a secondary step, have a *different* AI check the first one's work .
 - *Summary:* The seminar concluded by stressing that AIs can be wrong, and users must stay critical and always fact-check the information.