

文教大学大学院 人間科学研究科

2025年度 I 期

人間科学専攻

一般入学選考試験問題

【英 語】

受 験 番 号	氏 名	得 点

問 題 文

次の英文は、Stanislas Dehaene *How We Learn—The New Science of Education and the Brain* (Penguin Books, 2020) の Chapter2 ‘*Why Our Brain Learns Better Than Current Machines*’からの抜粋です。次の英文を読み、以下の問いに答えなさい。なお、解答は解答用紙に記入しなさい。

Characteristic of the human species is a relentless search for abstract rules, high-level conclusion that are extracted from a specific situation and subsequently tested on new observations. (a) Attempting to formulate such abstract laws can be an extraordinary powerful learning strategy, since the most abstract laws are precisely those that apply to the greatest number of observations. Finding the appropriate law or logical rule that accounts for all available data is the ultimate means to massively accelerate learning—and the human brain is exceedingly good at this game.

Let us consider an example. Imagine that I show you a dozen opaque boxes filled with balls of different colors. I select a box at random, one from which I have never drawn anything out before. I plunge my hand into it, and I draw a green ball. Can you deduce anything about the contents of the box? What color will the next ball be?

The first answer that probably comes to mind is: I have no idea—you gave me practically no information; how could I know the color of the next ball? Yes, but ... imagine that, in the past, I drew some balls from the other boxes and you noticed the following rule: in a given box, all balls are always the same color. The problem becomes trivial. When I show you a new box, you need only to draw a single green ball to deduce that all the other balls will be this color. With this general rule in mind, it becomes possible to learn in a single trial.

(b) This example illustrates how higher-order knowledge, formulated at what is often called the “meta” level, can guide a whole set of lower-level observations. The abstract meta-rule that “in a given box, all the balls are the same color,” once learned, massively accelerates learning. Of course, it may also turn out to be false. You will then be massively surprised (or should I say “meta-surprised”) if the tenth box you explore contains balls of all colors. In this case, you would have to revise your mental model and question the assumption that all boxes are similar. (c) Perhaps you would propose an even higher-level hypothesis, a meta-meta-hypothesis—for instance, you may suppose the boxes come in two kinds, single-colored and multicolored, in which case you would need at least two draws per box before concluding anything. In any case, formulating a hierarchy of abstract rules would save you valuable learning time.

(d) Learning, in this sense, therefore means managing an internal hierarchy of rules and trying to infer, as soon as possible, the most general ones that summarize a whole series of

observations. The human brain seems to apply this hierarchical principle from childhood. Take a two- or three-year-old child walking in a garden and learning a new word from his or her parents, say, *butterfly*. Often, it is enough for the child to hear the word once or twice, and voila! its meaning is memorized. Such a learning speed is amazing. It surpasses every known artificial intelligence system to date. Why is the problem difficult? Because every utterance of every word does not fully constrain its meaning. (c) The word *butterfly* is typically uttered as the child is immersed in a complex scene, full of flowers, trees, toys, and people; all of these are potential candidates for the meaning of the word—not to mention the less obvious meanings: every moment we live is full of sounds, smells, movements, actions, but also abstract properties. For all we know, butterfly could mean color, sky, move, or symmetry. The existence of abstract words makes this *think, believe, no, freedom, and death*, if the referents cannot be perceived or experienced? How do they understand what “I” means, when each time they hear it, the speakers are talking about ... themselves?!

The fast learning of abstract words is as incompatible with naïve views of word learning as Pavlovian conditioning*¹ or Skinnerian*² association. Neural networks that simply try to correlate inputs with outputs and images with words, typically require thousands of trials before they begin to understand that the word *butterfly* refers to that colored insect, there, in the corner of the image ... and such a shallow correlation of words with pictures will never discover the meanings of words without a fixed reference, such as *we, always, or smell*.

(出典：Stanislas Dehaene *How We Learn—The New Science of Education and the Brain*, Penguin Books、2020、pp.35-36)

*¹ Pavlovian conditioning … パブロフ型条件付け。ロシアの生理学者パブロフが行った、犬に餌を与える際に音を鳴らして生理反応をみる実験に基づく。一定の刺激を繰り返し与え続けることで、次第に生理的な条件反射が起こるようになるというもの。

*² Skinnerian … スキナー及びスキナーの行動主義心理学に関する、という形容詞。スキナーは「オペラント条件付け」などを唱えた 20 世紀アメリカの心理学者。

問 1 下線部 (a) ～ (e) を日本語に訳しなさい。

問 2 問題文冒頭にあるとおり、抜粋した英文の Chapter タイトルは '*Why Our Brain Learns Better Than Current Machines*' です。本文の内容を踏まえながら、今日における人間と機械との関係について、あなたの考えを日本語で述べなさい。

【「英語」の問題は以上です】