

DOCTORAL THESIS

Neural Representation of Anticipation Involved in Decision Making

Author:

Yumi SHIKAUCHI

Supervisor:

Dr. Shin ISHII

KYOTO UNIVERSITY

Abstract

Graduate School of Informatics

Department of Systems Science

Neural Representation of Anticipation Involved in Decision Making

by Yumi SHIKAUCHI

In 1971, O'Keefe and Dostrovsky reported that responses of cells in the dorsal hippocampus of rats to restraining tactile stimuli represent information involved in place identification. Neuroscientists have been studying the hippocampal place systems in rodents, primates and human until today. One of the most curious finding therein is that the hippocampal place system is just like a car navigation system: it represents not only a current position, but also possible future paths. However, how animals and humans detect/expect special information has not been well elucidated. We have no bird's-eye view indubitably, and have to connect external, partially observable visual cues to the map maintained in the brain for knowing the topography around us. Resolution of uncertainty due to the partial observability is an unexplained function of the brain. Moreover, sensory inputs (e.g., a view of scenery in vision) are often quite different from the map-like representation with a well-organized coordinate system. Thus, position detection should require unknown intermediate representation between the map-like topographic system and each scene view.

In this thesis, I describe the neural encoding mechanism that translates an upcoming scene view into its anticipating brain activity. To examine how to cope with the uncertainty resolution, I performed statistical modeling of human behaviors when performing a partially observable maze navigation task; in particular, I developed a hidden Markov model (HMM), which incorporates inference of a hidden variable in the environment and switching between exploration and exploitation based on the inference. The HMM-based model well reproduced the human behaviors, suggesting the human subjects actually performed exploration and exploitation to

effectively adapt to this uncertain environment. To determine how the brain performs view expectation during spatial navigation, I next applied a multiple parallel decoding technique to functional magnetic resonance imaging (fMRI) when human participants performed scene choice tasks in learned maze navigation environments. I could decode participants' view expectation from fMRI signals in parietal and medial prefrontal cortices, whereas activity patterns in occipital cortex represented various types of external cues. The decoder's output reflected participants' expectations even when they were wrong, corresponding to subjective beliefs which were opposed to objective reality. Next, I explored appropriate encoding models that well reproduce anticipating fMRI activities in a data-driven manner during maze navigation. The identified encoding models were noise-tolerant, because they incorporated the characteristics of the task environment. I also found that the encoding models could predict fMRI activities in the inferior parietal gyrus and precuneus, and that details of anticipated scenes were locally represented in the superior prefrontal gyrus. Furthermore, a decoder associated with the data-driven encoding models accurately predicted future scene views in both passive and active navigational environments.

In summary, my experimental studies indicated that humans anticipate a forthcoming scene during navigation, so that the neural representations of scene anticipation are robustly and effectively encoded in the parietal and medial prefrontal regions. In the decoding study, I demonstrated decoding of the participants' inaccurate prior belief of future observation. Additionally, my data-driven analysis showed that brain-activity-driven encoders represent scene views in a noise-tolerant manner and that the complementary decoder accurately predicts the future scene views. My findings through this experimental and analytical studies would be able to contribute to the future development of a new type of tool that allows people to communicate with non-linguistic modality.

Acknowledgements

First and foremost I would like to thanks to my adviser Professor Dr. Shin Ishii, you have been a tremendous mentor for me. I would like to thank you for encouraging my research and for allowing me to grow as a research scientist. He has taught me, both consciously and unconsciously, how good research is done. I appreciate all his contributions of time, ideas, and funding to make my Ph.D. experience productive and stimulating. Your advice on both research as well as on my career have been priceless. I would also like to thank my committee members, Professor Tetsuya Matsuda and Professor Manabu Kano for comments that greatly improved my Ph.D. thesis.

I would like to thank all of the Ishii Lab members. Dozens of people have helped and taught me immensely at the Ishii Lab. Thanks to Dr. Shigeyuki Oba and Dr. Shin-ichi Maeda for sharing their knowledge. Thanks to Dr. Hidetoshi Urakubo, Dr. Henrik Skibbe, Dr. Ken Nakae, Dr. Masanori Koyama and Dr. Naoto Yukinawa for giving valuable advice. My long PhD. student days were supported by collaborators: Mr. Masahiro Adomi, Mr. Yoshiki Azuma, Mr. Sotetsu Koyamada, Mr. Takuya Fuchigami and Mr. Yasuyuki Hamada. Dr. Naoki Honda, Dr. Yohei Kondo, Dr. Motoori Kon, Dr. Hiroshi Morioka, Dr. Kourosh Meshgi, Dr. Yohei Murakami, Dr. Takashi Kawase, Dr. Tsuyoshi Ueno, Mr. Ayumu Yamashita, Ms. Yuzhe Li , Mr. Masayuki Kouno and Mr. Satoshi Kambara have been actively interested in my work. I would especially like to thank lab secretaries. All of you have been there to support me by your special ability of clerical business.

My sincere thanks also goes to Dr. Motoaki Kawanabe, Dr. Haruo Hosoya, Dr. Jun-ichiro Hirayama, Dr. Takeshi Ogawa, Dr. Hiroki Moriya, Dr. Atsunori Kanemura, Mr. Satoshi Morimoto, and lab secretaries who provided me an opportunity to join their team as intern, and who gave access to the laboratory and research facilities. Without they precious support it would not be possible to conduct this research. Dr. Masaki Maruyama gives me kindly supports about ethical procedures for brain measurement. Also, thanks Prof. Ben Seymour for proofreading my journal paper. I want to thank Dr. Yutaka Nakamura. He supports my experiments. Dr. Takashi Takenouchi gives me kindly supports about data analyses.

Thanks to Dr. Michio Nomura and members of the Nomura lab for precious discussion and chats.

Thanks to Dr. Keiichi Kitajo and members of his lab for their help and chats.

Special thanks to my family: my mother-in law, father-in-law, my mother, and father. I would also like to thank to my beloved sisters and brother. Thank you for supporting me for everything. At the end I would like express appreciation to my beloved husband Dr. Manabu Shikauchi, who was always by my side.

Yumi Shikauchi

2016

Contents

Abstract	iii
Acknowledgements	v
1 Motivation for studying anticipation	1
1.1 Anticipation: focus of interest	1
1.2 Anticipation and intelligence	1
1.3 Anticipation and decision making	2
1.4 Anticipation and internal world	4
2 Existing study of anticipation in the brain	7
2.1 Reward expectation	7
2.1.1 Model-free system in dorsolateral striatum	7
2.1.2 Model-based system in prefrontal cortex	8
2.1.3 Integrative decision making process	9
2.2 Hippocampal place system	10
2.2.1 Mechanism of place coding	10
2.2.2 Replay	10
2.2.3 Action planning	11
2.3 Research direction	11
2.3.1 Research question	11
2.3.2 Measuring methods	12
Psychological experiment	12
Human brain activity	12
2.3.3 Analysis methods	13
Hidden Markov model	13
Multi-voxel pattern analysis	14
Encoding and decoding	14
Support vector machine	16

	Sparse logistic regression	17
3	View anticipation during spatial navigation	19
3.1	Experimental setting	19
3.1.1	Participants	19
3.1.2	Behavioral task	19
3.2	Analysis methods	21
3.2.1	Behavioral analysis	21
3.2.2	Computer agents	22
	Full-observation agent	22
	Stroll agent	22
	Optimal agent	23
3.2.3	Assumptions on human model	24
	Action selection probability	25
	Initial belief	28
	Estimation and prediction	28
3.3	Results	30
3.3.1	Parameter Estimation	30
3.3.2	Action prediction	30
3.3.3	Belief state estimation	31
3.4	Discussion	33
4	Decoding View Anticipation	35
4.1	Experimental setting	35
4.1.1	Participants	35
4.1.2	fMRI Data acquisition	35
4.1.3	Behavioral task	36
4.2	Analysis methods	39
4.2.1	Preprocessing for fMRI data	39
4.2.2	Regions of interest	39
4.2.3	Labeling of fMRI data	40
4.2.4	Decoding analysis overview	41
	Comparison between status-specific decoders	41
	Evaluation of the decoders for the upcoming scene view	43
4.2.5	Reconstruction methods	43

4.3	Results	44
4.3.1	Behavioral results	44
4.3.2	Decoding performance	45
	Time-course decoding analysis	45
	Upcoming scene view decoding from delay-period brain activity	47
4.3.3	Map reconstruction	49
4.3.4	Reconstruction and human behavior	50
4.4	Discussion	50
4.4.1	Status-specific decoding analyses	50
4.4.2	Decoding of expectation in decision making	53
5	Encoding View Anticipation	55
5.1	Experimental setting	55
5.1.1	Participants	55
5.1.2	fMRI Data acquisition	55
5.1.3	Behavioral task	56
	Pretraining	56
	Scene choice (SC) task	56
	Motion decision (MD) task	58
5.2	Anaysis methods	59
5.2.1	Behavioral analysis	59
5.2.2	fMRI data preprocessing	59
5.2.3	Encoding models	60
	Whole-scene encoding model.	60
	View-part encoding model.	60
	Full encoding model.	60
5.2.4	General linear model analysis	61
5.2.5	Data-driven analysis	61
	Predictability of encoding models	61
	Appearance frequency matrix of scene view	62
	Robustness of encoding models	63
	Characterizing information gains	63
5.2.6	Decoding analysis method	65
	Training of scene anticipation classifiers	65

Scene reconstruction	66
Time-shift Hamming distance	66
5.3 Results	67
5.3.1 Behavioral results for passive and active navigation	67
5.3.2 Comparison between encoding models	68
5.3.3 Scene anticipation in frontal and parietal cortices	75
5.3.4 Decoding results	77
5.4 Discussion	81
6 Conclusion and Future Directions	89
6.1 Functional network for anticipation	89
6.2 Neural representation for anticipation	90
6.3 Challenges for the future	91
Bibliography	93

List of Figures

1.1	Observations and time/space	2
1.2	Reinforcement learning	4
2.1	Analysis overview	9
2.2	Graphical model of cognitive state	13
2.3	Illustration of multi-voxel pattern analysis	15
3.1	A single block in the maze navigation game	21
3.2	Histogram of additional steps needed by human participants, in comparison to the full-observation agent	24
3.3	An HMM for human information processing during navigation behaviors	26
3.4	Navigation behaviors of a human participant and computer agents .	32
4.1	An example scene choice trial	38
4.2	Analysis overview	42
4.3	Behavioral results	46
4.4	Decoding results	48
4.5	Analysis overview	51
5.1	Experimental procedure and encoding scheme	57
5.2	Characteristics of encoding models	69
5.3	Averaged prediction accuracy over ten folds in the 10-fold CV in each individual	71
5.4	Minimum encoding model, participants 2-6	72
5.5	Minimum encoding model, participant 7	73
5.6	Effects of bit-inversion errors in the data-unrelated minimum encoding model	75
5.7	Brain regions involved in scene prediction	76
5.8	Individual results of spatial distributions of the information gain index	78

5.9 Decoding results 79

5.10 Decoding accuracy for participants 2-7 80

5.11 Sequential decoding in two MD blocks 82

5.12 Action history and TSHD during the motion decision task 83

6.1 Diagram of anticipation network 90

List of Tables

3.1	Three computer agents and a model of human participants	22
3.2	Results of parameter estimation for the HMM	30
3.3	Action concordance rate	31
3.4	Action concordance rate in the decision trial	31
3.5	Confidence concordance rate	33
4.1	List of six regions of interest	40

List of Abbreviations

1R	One-versus-the-Rest
2D	2-Dimensional
3D	3-Dimensional
BOLD	Blood-Oxygen Level Dependent
CV	Cross-Validation
dPFC	dorsal Prefrontal Cortex
ECOC	Error Correcting Output Coding
EEG	ElectroEncephaloGraphy
EPI	Echo-Planar Image
fMRI	functional Magnetic Resonance Imaging
GLM	General Linear Model
HC	Hippocampal Cortex
HMM	Hidden Markov Model
IGI	Information Gain Index
IPG	Inferior Parietal Gyrus
LOMO	Leave-One-Map-Out
LOSO	Leave-One-Session-Out
MD	Motion Decision
MEG	MagnetoEncephaloGraphy
ML	Maximum Likelihood
mPFC	medial PreFrontal Cortex
MVPA	Multi-Voxel Pattern Analysis
OC	Occipital Cortex
PC	Precuneus
paraHC	paraHippocampal Cortex
RO	Rolandic Opeculum
ROI	Region Of Interest
RT	Reaction Time

SC	Scene Choice
SFG	Superior preFrontal Gyrus
SLR	Sparse Logistic Regression
sPC	superior Parietal Cortex
SVD	Singular Value Decomposition
SVM	Support Vector Machine
TD	Temporal Difference
TP	Temporal Pole
TSHD	Time-Shifted Hamming Distance

Chapter 1

Motivation for studying anticipation

We can recall vividly what happened in the past, possibly as images, even without external stimuli. Similar to or freer than reliving, we can envision what will happen in the future. In this chapter, I define the notion of future anticipation with some scenarios in our daily life, and describe how the processing of anticipation remains as a mystery of the brain function.

1.1 Anticipation: focus of interest

You can imagine the scene behind the door before opening it. You know what are in your drawer even when it is closed. Like these we often have knowledge about the outside of our sight. In other words, many things that we have in mind are not observed directly due to occlusion, remoteness or having not yet happened. Here, I classify our observations, which are received through our sensory modalities (sight, hearing, smell, taste, and touch), into three categories based on the difference on the time axis: past observations, current observations, and future observations (Fig. 1.1). When we predict the future observations, they are called anticipation.

1.2 Anticipation and intelligence

We predict future observations by integrating the observations available in the past and at present. For example, if you hear the sound of water falling on the roof then you think it is raining outside. If you placed your scissors in the drawer yesterday, then you are sure that you find it there. We instinctively believe that there will be tomorrow that is similar to but slightly different from today. It is natural that we

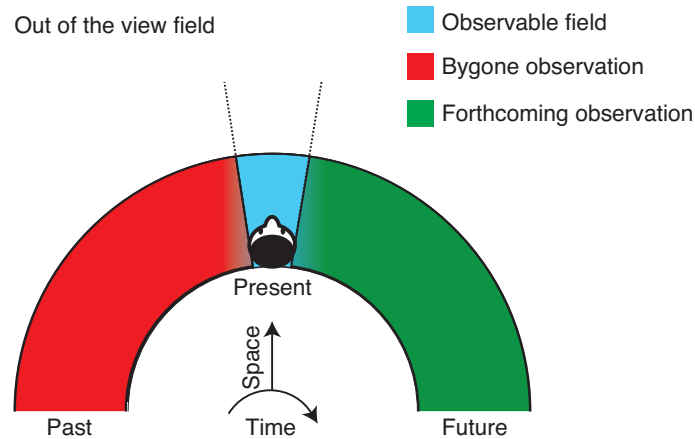


FIGURE 1.1: Observations and time/space. The visual field is limited in terms of both of time and space. Anticipation is defined as mental image of forthcoming observations (green area) throughout this thesis.

believe that the real world is subject to the laws of physics, such as to be governed by continuity because of the law of inertia.

Our prediction of future events is based on the continuity from the past to the present. The continuity of the world is implicitly recognized by ourselves. Our anticipation, which is performed in our everyday life with no difficulty, is a typical intellectual activity; a wide range of knowledge is necessary for performing anticipation. Laughter provides some evidence on intellectual aspects of anticipation. As it is said that only human can laugh, laughter is one expression of intelligence. We observe that incongruity between anticipation and actual observation is one important factor for causing laughter (Gervais and Wilson, 2005). Although the major elicitor of laughter is non-serious social incongruity, incongruity in serious contexts can induce an opposite emotion, for example, crying by human infants. Accordingly, anticipation of future observations is one of the higher-order mental activities that arise in our daily life.

1.3 Anticipation and decision making

In the previous section, I described the emotional reactions associated with anticipation, such as laughing and crying. Besides, anticipation plays an important role in behavioral reactions, that is, decision making. This section explains links to decision making based on behavioral experiments and computational interpretations.

When a specific action is accompanied by positive outcome, then an animal comes to perform that action repeatedly. For example, Skinner (1938) examined the behaviors of rats placed in a small box with a lever, which is known as the Skinner's box. A naive rat seldom pressed the lever. However, after the rat was fed after it happened to press the lever several times, the frequency of the lever press gradually increased. This result suggests that the rat learned the lever press that was associated with the preferable outcome, i.e., food. To establish this association, the rat needs to predict upcoming food reward (result) depending on its preceding action (cause). After many assiduous efforts, it has been convincingly argued that such behavioral changes induced by the outcome, instrumental conditioning, occur in a consistent manner over many species of animals (Skinner, 1938).

The theory of conditioning, classical and instrumental, has been confirmed by huge number of psychological experiments, then impacted not only on ethology but also other fields, like economics and artificial intelligence. Barto, Sutton, and Anderson (1983) presented an idea of reinforcement learning, which was primarily an extension of the Rescorla-Wagner model (Rescorla and Wagner, 1972) for classical conditioning. Robots mimic the process of animal learning to make them adapt to its surrounding environment flexibly. Reinforcement learning is widely accepted as a comprehensive theoretical framework that covers classical conditioning, instrumental conditioning and their time-series extensions (Fig. 1.2). In reinforcement learning, an animal/robot (agent) selects an action (a) based on observation (o) from the environment. Then, the environment provides the agent with a reward (r), which is an assessment of the agent's action. The goal of the agent is to seek such an optimal policy that maximizes the temporal accumulation of the reward r . The expectation of the accumulated reward is called the value, and when we intend to show its dependence on the current state, it is called the value function. The agent thus acquires the optimal policy by trial and error, based on the identification of the value function. When there are multiple candidates of action, a fairly good policy can be given such to more frequently select an action with a larger value. The so-called soft-max policy is used as such an instance:

$$p(a) = \frac{\exp(\beta m_a)}{\sum_{a'=1}^{N_a} \exp(\beta m_{a'})}. \quad (1.1)$$

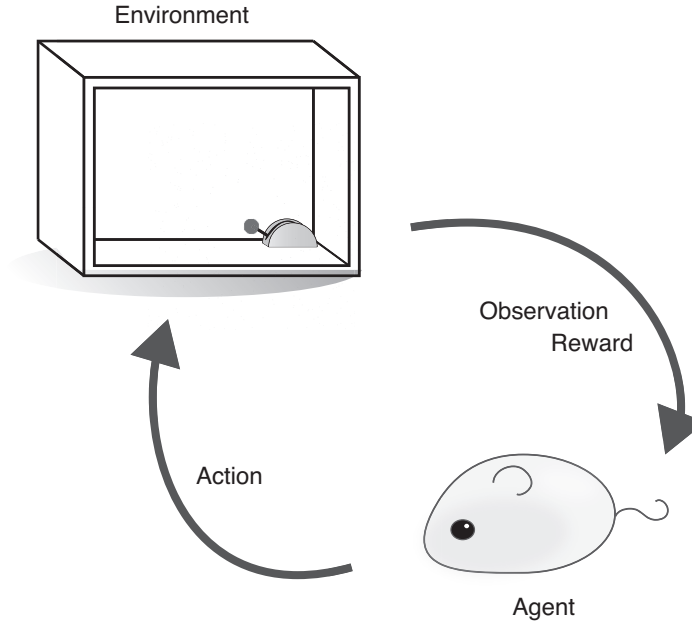


FIGURE 1.2: Reinforcement learning.

Here, m_a and N_a are the action value and the number of action candidates, respectively. Note that $\sum_{a'=1}^{N_a} p(a') = 1$. The parameter β sometimes called as an inverse temperature, controls the steepness of the soft-max function, and determines the variability of the agent's action. The action value m_a is adjusted by a stochastic approximation of the dynamic programming (Sutton, 1988), as follows:

$$m_a \rightarrow m_a + \epsilon(r_a - m_a). \quad (1.2)$$

Here, ϵ is a parameter known as the learning rate, which determines to what extent the newly acquired reward following action a (r_a) overrides the old action value. Note that this implementation is one example of reinforcement learning.

1.4 Anticipation and internal world

We, humans sometimes have to anticipate forthcoming situations in order to make appropriate decision, such as in a complicated environment that may include various kinds of uncertainty. A typical example can be seen when we navigate through a familiar city based on observations of scenes around us. I mentioned above the notion of instrumental conditioning in which the reward anticipation can be directly

associated with the current action. However, we sometimes face more intricate situations; for example, we may receive a reward after sequence of actions, rather than a specific single action, or get a reward with non-negligible time delay, such as salary paid monthly. Here, I show a computational framework to deal with such more complicated cases.

The formula (1.2) showed the relationship between action a and reward r at the current observation o . Here, I introduce another variable, state s , which may be a function of a , r , and/or observation o . This representation can describe some context-dependency, such as whether or not a specific event has been finished. The state includes every information necessary for making appropriate decision, including an initial state, the history of action, reward, and/or observation. Then, the goal of a reinforcement learning agent is to seek such an optimal policy that maximizes the total reward, given a current state. When a reinforcement learning agent is able to identify explicitly a model of state transition, which is called an internal model, it is said as performing model-based reinforcement learning. Although it is in many realistic cases difficult to identify the internal model in an efficient manner, the total reward acquired by the reinforcement learning agent can almost be the maximum after the agent identifies the internal model of the current environment.

When you are in your familiar city, you can go anywhere you like without any paper map, based on the environmental model in your mind, which is an instance of the internal model, and in this particular case, called a cognitive map. We can construct the cognitive map based on our empirical knowledge, and then know the current position in the city and available route from the current position to the destination. If we do not have such an internal model, it is impossible to plan possible route toward the destination.

Chapter 2

Existing study of anticipation in the brain

Neural correlates of future events have been studied in two major contexts: reward expectation is ubiquitous in the brain during value-guided decision making (Doll et al., 2015); and, during spatial navigation, hippocampal place cells depict potential future paths (Pfeiffer and Foster, 2013; Johnson and Redish, 2007). In this chapter, I review existing studies on these two lines and clarify my research motivations in relations with them.

2.1 Reward expectation

I here describe theoretical frameworks of model-free (Section 1.3) and model-based (Section 1.4) decision making. The brain is considered to contain these two different modalities of decision making in parallel.

2.1.1 Model-free system in dorsolateral striatum

Dopamine neurons originally attracted neuroscientists' attention due to possible involvement in Parkinson's disease. In 1986, Schultz (1986) found that monkeys' midbrain dopamine neurons responded when food reward was delivered to the mouth. Subsequently, he and his colleagues found that dopaminergic neurons in primates are activated by the reward-predicting stimulus after stimulus-reward conditioning (Schultz, Apicella, and Ljungberg, 1993). In 1997, they linked these neural representation and a theoretical concept, reward prediction error (Schultz,

Dayan, and Montague, 1997):

$$\delta_t = r_t + \gamma V_{t+1} - V_t. \quad (2.1)$$

Here, r_t is the reward at time t and V denotes the expectation of the sum of future rewards up to the end of the trial. When summing, each reward r is discounted depend on the temporal distance; rewards that arrive sooner are more important than rewards that arrive later. $0 \leq \gamma \leq 1$ is the discount factor. This δ_t is called the temporal difference (TD) error. The firing rate of dopamine neurons reflects the magnitude of reward prediction error and/or probability of error occurrence of all-or-none reward (Tobler, Fiorillo, and Schultz, 2005; Fiorillo et al., 2003), which could be responsible for prediction-based learning (Rescorla and Wagner, 1972; Pearce and Hall, 1980) and TD learning (Sutton, 1988). These studies suggested that dopamine neurons are closely involved in expectation of reward.

2.1.2 Model-based system in prefrontal cortex

In complicated environments, e.g., situations of sequential behaviors to realize future goals, appropriate decision making has to include planning based on an internal model. When utilizing or even constructing the internal model of current environment, the necessary information is often partially accessible, i.e., there is uncertainty. Then, gathering information to resolve the uncertainty is one of the essentials of applicable planning, as well as approaching rewards. Early lesion studies suggested a relationship between impairments in planning and damage in the frontal lobe (for a review, see Owen (1997)).

For examining exploratory decision, Daw et al. (2006) showed human behaviors and fMRI brain activities in an uncertain environment. Each human subject was requested to get as much as rewards (payments) by iteratively choosing one slot machine from four options at each time. Each slot machine had a time-variant payoff, whose time-series was unknown by the subjects. For maximizing the total amount of rewards, the subjects tried to know the best option based on trials and errors (exploratory behaviors). After they got to believe that a specific option is the best, they repeatedly chose it (exploitative behaviors). Bilateral frontopolar regions showed significantly increased activation accompanied by exploratory behaviors in comparison to exploitative behaviors. This study implied that decision making with

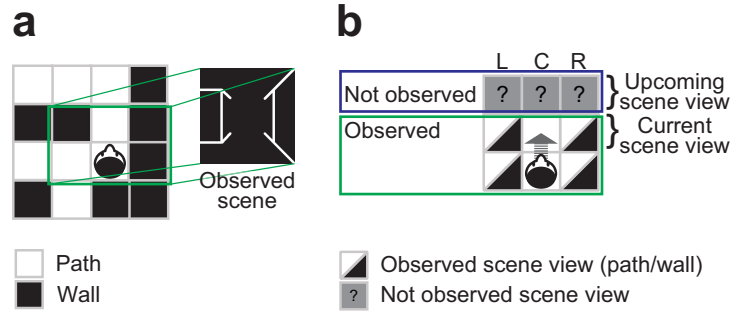


FIGURE 2.1: A conceptual scheme of the navigation environment. The environmental information consists of already seen and unseen squares. There are two types of squares: white paths and black walls. (a) An example of the maze environment and scene view. A green box indicates a field of view for this participant. (b) After moving forward, here from the bottom-center square to the middle-center square in the green box, this participant observes the wall statuses of top-left, top-center and top-right squares. L, left; C, center; R, right.

uncertainty involves switching between exploratory and exploitative behavioral modes.

Around the same time, Yoshida and Ishii (2006) reported neural processes involved in navigation in an uncertain environment. The subjects were requested to sequentially select appropriate actions for goal achievement in a virtual three dimensional (3D) maze (Figure 2.1a). In this navigation task, they were not informed of the initial position. Soon after the start of a navigation block, the subjects tried to identify the current position using available observation (3D scene views), which is a partial information of the current position, and the knowledge of maze; they were familiar with the maze structure due to pre-training sessions outside the MRI scanner. Namely, they reduced the number of position candidates (i.e., uncertainty of the position) based on their exploratory behaviors. This study concluded that the uncertainty is represented and resolved in the anterior prefrontal cortex.

2.1.3 Integrative decision making process

The brain makes decisions using these dual action choice systems; the model-free midbrain system and model-based prefrontal system. One of normative questions is then how the brain utilizes in an integrative fashion the multiple control systems for selecting appropriate actions. The model-free controller is directly motivated by

rewards, whereas the model-based controller needs to incorporate also the uncertainty of the environment. These two controllers may disagree in some cases because of the trade-off between, for example, certain small rewards and less certain but large rewards. This trade-off is known as ‘exploration-exploitation dilemma’. Daw, Niv, and Dayan (2005) showed a model of uncertainty-based controller competition, whose computer simulation well reproduced a rat’s reinforcement learning task. According to the model, when a rat confronted high uncertainty (i.e., soon after the initial trial), the model-based system dominated. After many trials, the model-free system prevailed, which was similar to habituation (Dickinson, 1985). This trade-off would be represented in human brain; when the model-based controller was dominant, the putamen blood-oxygen level dependent (BOLD) response being correlated with model-free prediction error decreased, but was increased when the model-free was dominant (Doll et al., 2015).

2.2 Hippocampal place system

2.2.1 Mechanism of place coding

In 1971, O’keefe and Dostrovsky reported that responses of cells in the dorsal hippocampus of rats to restraining tactile stimuli behave as a function of spatial orientation (O’Keefe and Dostrovsky, 1971). The hippocampal place system is now believed to be a well-established system for representing spatial information in the brain (Buzsáki and Moser, 2013). The hippocampal cortex and the parahippocampal cortex belong to this system. The place cells in the hippocampal cortex receive multiple grid cells’ input (Hafting et al., 2005; Fyhn et al., 2008, for a review, see Penner and Mizumori (2012)). The parahippocampal cortex consists of grid cells, which show firing specific to location. The grid cells are of diverse scales and directions. Thus, a place cell encodes a specific position by overlaying multiple outputs of the grid cells.

2.2.2 Replay

The place cells have long been known to be involved in spatial decision making such as navigation. Foster and Wilson (2006) provided a novel insight into the learning process of spatial information. They reported a specific activity in the awake rest;

the hippocampal place cells behave sequential replay, in which recent episodes of spatial experience are repeated as reverse reproduction. In the rat's brain, the reverse replay is more frequently observed in a new environment than in a familiar one. This study suggested that the reverse replay has a critical role in support of learning and memory consolidation during hippocampus-dependent tasks. Following this study, several neuroscientists have reported replay in the brain during not only the awake rest but also the sleep (Eschenko et al., 2008; Epstein, 2008; Girardeau et al., 2009). These studies suggested that the hippocampal place system realizes memory representation dependent on the order but rather independent of the period: the latter is called time-invariance below.

2.2.3 Action planning

This time-invariance could be attained for future, too. The place cells represent spatial locations along an absolute coordinate in terms of a spatial map and moreover, are known to depict potential future paths. When the rats are running on a T-shape maze, the hippocampal place cells represent current position as mentioned Subsection 2.2.1. Johnson and Redish (2007) found that the spatial activity is ready to estimate the rat's current position. The rat's hippocampus performs in terms of its sequential activity encoding trajectories in a start position on an open-field (Pfeiffer and Foster, 2013). These activities implied the rat's future trajectory for approaching an reward position. This future-trajectory-dependent firing in the hippocampal cortex is represented by a wide circuit including structures involved in the evaluation and selection of actions: the source of this activity is the medial prefrontal cortex projected via the midline thalamic nucleus reuniens (Ito et al., 2015).

2.3 Research direction

2.3.1 Research question

Under the existing studies above, in this thesis, I target future anticipation in human decision making. In the scheme of model-based decision making, how the human brain encodes anticipation? To answer this question, I measured human brain activity during spatial navigation games, and analyzed the measured brain activity and behaviors.

2.3.2 Measuring methods

Psychological experiment

To examine human cognitive states, I recorded subjects' decisions using a button pressing during the subjects were performing spatial navigation games. Each button pressing provides two kinds of information: an internal state/action and a reaction time (RT). When we are interested in the internal state, each subject was requested to report his/her cognitive state using the button, so that the button-press leads to the evaluation of subjective perception and/or cognition. The RT represents the latency from a stimulus onset to a subject's response. Since the subjects were requested to press a button as soon as possible, RT could be seen as the measure of cognitive load: a higher load needs a longer RT.

Human brain activity

For measuring brain activities of healthy human, a non-invasive method is the sole choice. There are three major modalities for measuring human brain activity without any surgical operations or radiation exposure, which have been widely used in neuroscience studies: electroencephalography (EEG), magnetoencephalography (MEG) and functional magnetic resonance imaging (fMRI). EEG is a modality for monitoring electrical activity of the brain. When measuring human brain activity, multiple electrodes are placed on the scalp. Each electrode detects voltage fluctuations resulting from ionic current within multiple neurons over the cerebrospinal fluid, skull and scalp. Similar to EEG, MEG detects brain activity based on electric currents in neurites. MEG measure changes in magnetic fields by means of high-sensitivity SQUID coils, caused by electrical activity of the brain. These two functional neuroimaging modalities provide fine temporal resolution. However, EEG and MEG serve poor spatial resolution because the brain signals measured at each sensor position are mixtures of original signals produced by different brain regions. The decomposition of measured signals into original signals constitute a difficult inverse problem. On the other hand, fMRI provides fine spatial resolution with an active sensing technique. In fMRI, a particular brain slice is excited by a magnetic field gradient and radio frequency pulse (selective irradiation). Magnetization of oxygen-rich blood is different from that of oxygen-poor blood (blood-oxygen-level

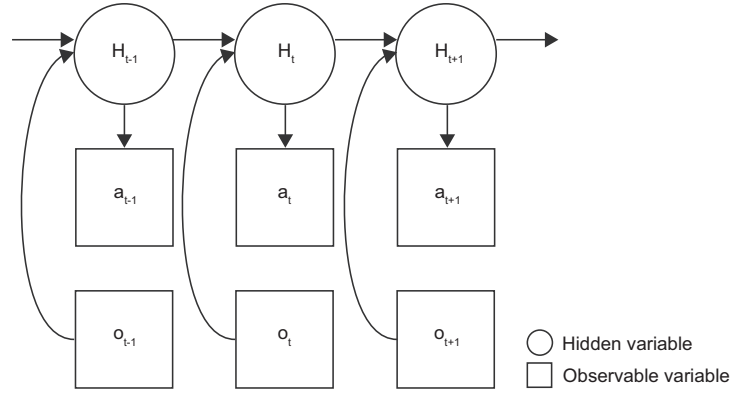


FIGURE 2.2: Graphical model of cognitive state. H , a and o denote a hidden state, action and observation (e.g., stimulus or reward), respectively. Subscripts indicate a time variable.

dependent (BOLD) contrast) (Huettel, Song, and McCarthy, 2008). Thus, fMRI detects the BOLD signal which is caused by local increase in the blood flow, which is in turn caused by the neurons' consumption of oxygen during their activation. Among these different modalities, fMRI has the highest spatial resolution so as to obtain the most accurate activity-related signals, being unaffected by noises coming from, for example, the outside of cortical regions. In this study, I measured functional brain activity as time series of three-dimensional images taken by fMRI.

2.3.3 Analysis methods

Hidden Markov model

An action (e.g., a selected button type and reaction time in psychological experiments) is one of indications of the subject's cognitive state. A stimulus (e.g., a visual image) modified the subject's cognitive state. An experimenter inferred an unobservable cognitive state of the subject from these observable information. Namely, I hypothesize that the subject decides his/her action based on his/her belief, which is dependent on the present stimulus (Fig. 2.2). The belief may change over time with a Markov property. This restriction of the temporal change of belief is appropriate for describing human cognitive states. To model the cognitive state from the time sequence of actions and stimuli, I used a hidden Markov model (HMM).

Multi-voxel pattern analysis

The traditional analyses for fMRI brain activity were often based on a subtraction method, which usually focuses on individual brain voxels. The main question in the subtraction method is whether each voxel activity is increased or unchanged while a subject performs a target cognition, like in cognitive tasks, compared to a base line condition. This approach has been tremendously productive, but has several limitations due to short of information about spatial patterns (Norman et al., 2006). By applying multivariate methods to fMRI data, it is possible to characterize functional relationships between different brain regions or sub-region voxels. This approach is known as multi-voxel pattern analysis (MVPA) (Fig. 2.3). The main benefits of the MVPA approach are as follows;

- By avoiding the signal-loss issues, the MVPA approach extracts the signal that is responded in the pattern in high sensitivity.
- The subtraction analysis is a condition (a few minutes)-based approach, while the MVPA approach detects the presence/absence of cognitive states based on individual trials (a few seconds).
- The MVPA approach may reveal the brain's way to combine multiple cognitive states into activity patterns with a very large capacity.

These benefits would lead to a solution to the most fundamental question in cognitive neuroscience, how to represent the external world by means of internal neural space, i.e., encoding. The MVPA approach have developed through visual sensory areas (Haxby et al., 2001; Kamitani and Tong, 2005; Kay et al., 2008). Several studies applied this method to auditory areas (Formisano et al., 2008; Meyer et al., 2010). However, it remains a challenge to decode brain activities in higher brain areas.

Encoding and decoding

Several MVPA studies have focused on decoding properties of the subject's cognitive state. fMRI-based decoding studies have recently provided new insights into the encoding schemes of the human cortex (Miyawaki et al., 2008; Stansbury, Naselaris, and Gallant, 2013; Nishimoto et al., 2011; Horikawa et al., 2013). For the brain, encoding is the process of translating perceived stimuli into nervous activity, while decoding is its complement, to reproduce environmental stimuli based

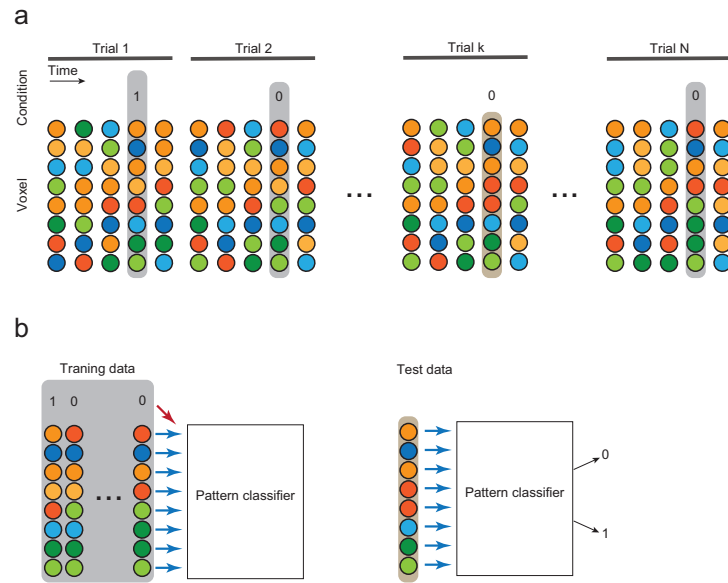


FIGURE 2.3: Illustration of multi-voxel pattern analysis. (a) An example of fMRI experiment setting and fMRI brain activities. The experiment consists of N trials with two types of task condition; 0 or 1. Each trial, which consists of several time points, belongs to either of two types of task condition. fMRI gives an experimenter a time series of brain activity pattern. Each element is called a voxel (colored circle), which represents a BOLD signal in a particular region. (b) MVPA consists of two phases; training phase and test phase. In the training phase, a classifier function is trained to map between brain patterns and task conditions (left). In the test phase, the classifier is applied to a novel pattern to predict the corresponding task condition.

on the measured brain activity, being comparable to the processing of the system downstream of central nervous system. Due to the indeterminacy in the general neural coding, the mimicking this decoding (termed simply decoding, hereafter), even with the help of full utilization of machine learning techniques, would show limited performance without knowledge of the complementary encoding scheme. This means, on the other hand, that decoding performance is a barometer of the reliability of our knowledge of the paired encoding scheme (Naselaris et al., 2011). Thus, in the study, shown in Chapter 4, I demonstrated neural decoding from fMRI brain activity pattern. Moreover, shown in Chapter 5, I selected an encoding model in a data-driven manner, and evaluated this model by a decoding accuracy.

Support vector machine

When analyzing fMRI multi-voxel pattern in chapter 5, I used a support vector machine (SVM). An SVM classifier finds the best hyperplane that separates training data points $\mathbf{x}$ of one class $t = -1$ from those of the other class $t = 1$, using linear models of the form

$$y^{SVM}(\mathbf{x}) = \mathbf{w}^T \phi(\mathbf{x}) + b \quad (2.2)$$

where $\phi(\mathbf{x})$ and b denote a fixed (and may be underlying) hyperplane transformation function and a bias parameter. The sign of $y^{SVM}(\mathbf{x}_n)$ corresponds to the class prediction of a new test point $\mathbf{x}_n$; i.e., $y^{SVM}(\mathbf{x}_n) < 0$ indicates the SVM predicts as $t_n = -1$ and $y^{SVM}(\mathbf{x}_n) > 0$ indicates its opposite prediction $t_n = 1$. The decision boundary is chosen to be the one for which the margin, perpendicular distance between the closest data points to the boundary, is maximized under the constraint as follows:

$$\arg \min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2. \quad (2.3)$$

The training of the SVM was performed so as to minimize the cost function:

$$L(\mathbf{w}, b, \mathbf{a}) = \frac{1}{2} \|\mathbf{w}\|^2 - \sum_{n=1}^N a_n \{t_n (\mathbf{w}^T \phi(\mathbf{x}_n) + b)\} \quad (2.4)$$

where $a_n \geq 0$ is the Lagrange multiplier. If the set of training data points are linear separable, the trained SVM will give perfect separation; however, such naive classifier can be poor in generalization. Moreover, in many cases, the data space itself may not be linear separable. Thus, we need to modify the SVM so as to allow some of the training points to be misclassified. To do this, the original cost function, equation 2.3, is replaced with

$$\arg \min_{\mathbf{w}, b} C \sum_{n=1}^N \xi_n + \frac{1}{2} \|\mathbf{w}\|^2 \quad (2.5)$$

where $C > 0$ is a regularization parameter, which controls how much the slack variable ξ is allowed to take a large positive value. The corresponding Lagrangian is given by

$$L(\mathbf{w}, b, \mathbf{a}) = \frac{1}{2} \|\mathbf{w}\|^2 - C \sum_{n=1}^N \xi - \sum_{n=1}^N a_n t_n y^{SVM}(\mathbf{x}_n) - 1 + \xi_n - \sum_{n=1}^N \mu \xi_n \quad (2.6)$$

where $a_n \geq 0$ and $\mu \geq 0$ are Lagrange multipliers.

Sparse logistic regression

In chapter 4, I applied another sparse estimation method, called sparse logistic regression (SLR), which is a logistic regression with a prior called an automatic relevance determination (Yamashita et al., 2008). SLR assumes a Gaussian prior distribution with a zero mean vector and a spherical covariance matrix instead of L2 regularization (Eq. 2.4):

$$y^{SLR}(\mathbf{x}_n) = \frac{1}{1 - \exp(-(\mathbf{w}^T \mathbf{x}_n + b))} \quad (2.7)$$

$$P(\mathbf{w}, b | \alpha) = N(0, \alpha^{-1} I_D). \quad (2.8)$$

Here, I_D is the identity matrix of size $D \times D$. D is the number of voxels.

To find the parameter vector $\mathbf{w}$ and the bias b , the following log-likelihood function is maximized.

$$L(\mathbf{w}, b) = \sum_{n=1}^N [t_n \log y^{SLR}(\mathbf{x}_n) + (1 - t_n) \log(1 - y^{SLR}(\mathbf{x}_n))]. \quad (2.9)$$

This maximization can be done by the Newton method, because the gradient and Hessian matrix of $L(\mathbf{w}, b)$ can be obtained in the closed forms.

The parameter vector $\mathbf{w}$ is sparsely estimated by means of the automatic relevance determination (Eq. 2.8). Moreover, in SLR, voxels are selected based on contribution of cross validation (CV) accuracy inner-training data by a selection counting value $SC(d)$:

$$SC(d) = \sum_{k:p(k)>p_{chance}}^K I(\hat{\mathbf{w}}_d(k) \neq 0) \times p(k) \quad d = 1, \dots, D. \quad (2.10)$$

Here, p_{chance} is a chance level with classification performance. $\hat{\mathbf{w}}(k)$ and $p(k)$ denote the estimated parameter vector and classification accuracy resulting from the k th CV. Note that the selection counting values are calculated using the training data set, then an additional data set, like a test data set, is not required.

Chapter 3

View anticipation during spatial navigation

In this chapter, I report behaviors in decision making and their modeling in partially observable environment. Several studies reported switching behaviors of humans between exploration and exploitation; however, computational details underlying it have been unexplored. To examine this issue, I used a 3D navigation game to which I applied a computational behavioral analysis based on hidden Markov model (HMM).

3.1 Experimental setting

3.1.1 Participants

Eight healthy participants (2 females, 6 males; aged 22–26 years, right-handed and with normal or corrected-to-normal vision) took part in this experiment. Each participant gave written informed consent, and all experiments were approved by the Ethics Committee of the Advanced Telecommunications Research Institute International, Japan. All methods were carried out in accordance to the approved guidelines. They practiced the free exploration task on the day before they performed the navigation task, in order to memorize the structure of the partially observable environment.

3.1.2 Behavioral task

In the experiment, human participants took part in a partially-observable maze navigation task, whose scheme is shown in Fig. 3.1. The participants were informed

of a goal position on the 2D maze, map at the beginning of a single block, and then tried to reach the instructed goal after as few trials as possible, based on partial 3D observations which gave clues to identify the current position on the maze. Notice that the participants were requested to reach an instructed goal from an uninstructed start position and orientation. Thus, the participants have to identify the current position before planning to achieve the goal.

Each participant performed three sessions each of which consisted of 150 trials of the navigation task. When a participant achieved a goal, it was the end of a single block, then another block started. The participants shared the same settings of the task; in the corresponding block, the start and goal positions were the same over the participants.

On the day before performing the maze navigation task, the participant had joined a free exploration task to well identify the maze structure and the correspondence between locations on the 2D maze map and 3D views which can be seen at those locations. The maze used in the free exploration task was common with the one in the navigation task done in the subsequent day, and I confirmed that they had sufficiently memorized the maze structure and also attained the mapping from 2D locations to 3D views through the free exploration task.

The maze I used (Fig. 3.1) consists of 11×11 squares: 58 vacant squares constituting paths and 63 wall squares. There is no dead-end, cross road, or wider (1×2 , 2×1 , etc.) square area. The observation (view) is 3D and partial; at each square on the maze, the participant may take either of four orientations, north-faced, east-faced, south-faced or west-faced. A single view presents the status, wall or vacant, of six squares (forward, forward-left, forward-right, left, right and the current position). There are in total seventeen possible views and all of them actually exist in the used maze. Each location on the maze shares the same view with some others, at least other five different locations on the maze; this maze design did not allow the participant to identify his/her position only from a single 3D view, which made the environment uncertain.

In a single trial in a block, after presenting a 3D view, the participant selected one of three possible actions, forward move, left turn and right turn, by a button press. This action selection should be completed within 2 seconds, which defines the maximum reaction time. If the forward move was selected, the location was changed by one square to the forward direction, but the other two actions only changed the

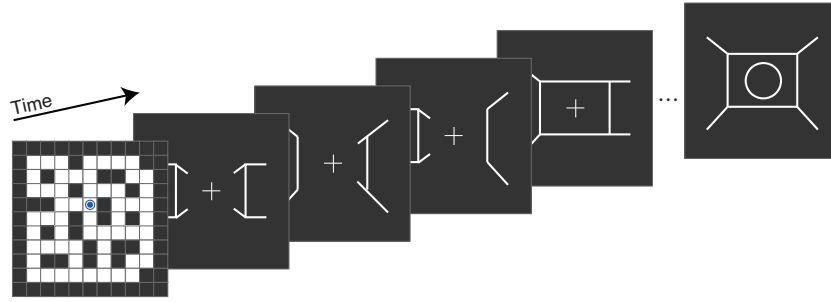


FIGURE 3.1: A single block in the maze navigation game. An 11×11 grid field was used. A blue circle on the map indicates a goal position. Each participant was requested to reach the goal from an uninstructed initial state based on sequentially observed 3D scene views. Each block continued until she/he achieved the goal arrival, which was informed by a white circle.

orientation with letting the participant stay at the same square. The participant reported his/her actions by her/his left hand; forward move by a middle finger, right turn by an index finger and left turn by a ring finger.

By using another hand, the participants were requested to report whether the actual observation was the same as what they had predicted; I assumed that the participant could predict the next view before it was presented. The participant was requested to first press the right-hand button to confirm their prediction and then the left-hand button to select their actions, in each trial. The report by the right hand was optional, but if the participant could not press the action button by the left hand within an allotted time (2 sec.), it caused a miss trial and he/she stayed at the same location with the same orientation.

3.2 Analysis methods

3.2.1 Behavioral analysis

All behavioral data analyses were performed with MATLAB (MathWorks) with the usage of the Statistics and Machine Learning toolbox. When searching for the shortest route in each block, I used Dijkstra's algorithm implemented as a custom program in MATLAB.

3.2.2 Computer agents

In order to evaluate human behaviors in this navigation task, I designed three computer agents; a full-observation agent, a stroll agent, and an optimal agent (Table 3.1). I simulated navigation processes of these computer agents in the same settings as those done by human participants (Subsection 3.1.2). Comparisons of the behaviors between the computer agents and human participants allowed me to quantitatively evaluate human decision making in this partially observable environment.

TABLE 3.1: Three computer agents and a model of human participants.

	Full-observation	Optimal	Human	Stroll
Memory	✓	✓	?	×
Initial state	✓	×	×	×

Full-observation agent

I first modeled a ‘Full-observation’ agent; this agent knows the perfect map topography and where it is on the map, and moreover, memorizes all the past observations. This agent takes the ideal, shortest route in all the trials in each block, so that the action selection probability is given by

$$p(a_t|s_t) = \begin{cases} \frac{1}{N_t^{shortest}} & \text{if } a_t \text{ is exploitative} \\ 0 & \text{otherwise} \end{cases} \quad (3.1)$$

where, s_t is the location (in the case of the full-observation agent, this location is the real one) and $N_t^{shortest}$ is the number of actions that constituted the shortest route, at the t -th trial.

Stroll agent

Second, I modeled an opposite of the full-observation agent, a ‘stroll’ agent. The stroll agent has no memory, but knows the effective action in this maze navigation environment; e.g., when a forward square is a wall, the forward movement to the

wall square is ineffective and a turn to left or right is effective. Based on such knowledge, I defined a model of action preference of the stroll agent as follows.

1. If effective, the agent prefers forward movement.
2. If there is a wall in the forward direction but no wall on both of left-hand side and right-hand side, the agent takes either of turn-to-left or turn-to-right with equal probability to see a vacant square in the new forward direction.
3. If the agent is on an L-shape corner, it takes the sole effective turn action.

Namely, this agent's action depends only on the current observation (denoted as o_t^*), not on the estimated location (denoted as s_t) or the real location (denoted as s_t^*). The action selection probability of this agent is given as follows.

$$p(a_t|o_t^*) = \begin{cases} \frac{k_{prefer}}{k_{prefer}N_t^{prefer} + (N_t^{eff} - N_t^{prefer})} & \text{if } a_t \text{ is the preferred action} \\ \frac{1}{k_{prefer}N_t^{prefer} + (N_t^{eff} - N_t^{prefer})} & \text{otherwise} \end{cases} \quad (3.2)$$

where, N_t^{eff} and N_t^{prefer} denote the numbers of effective actions and preferred actions, respectively, at the t -th trial. The parameter k_{prefer} defines the tendency to select the preferred action and was fixed to 0.90 (this agent almost selected the preferred action).

Optimal agent

Third, I modeled an 'optimal' agent; this agent memorizes all the past observations and takes an 'info-max' action which most reduces the possibilities of the agent's location on the maze if there is any position ambiguity, whereas a 'greedy' action to follow the shortest route to the goal if there is no ambiguity. Moreover, this agent is assumed to completely know the map topography. Then this agent selects actions similar to the stroll agent in the early part of each block, because the preferred action of the stroll agent is same as the 'info-max' action. When there is ambiguity, on the other hand, the action by the optimal agent was given by Equation 3.2 with k_{prefer} was fixed to 1. When there is no ambiguity, this agent selected an action according to Equation 3.1, similar to the full-observation agent.

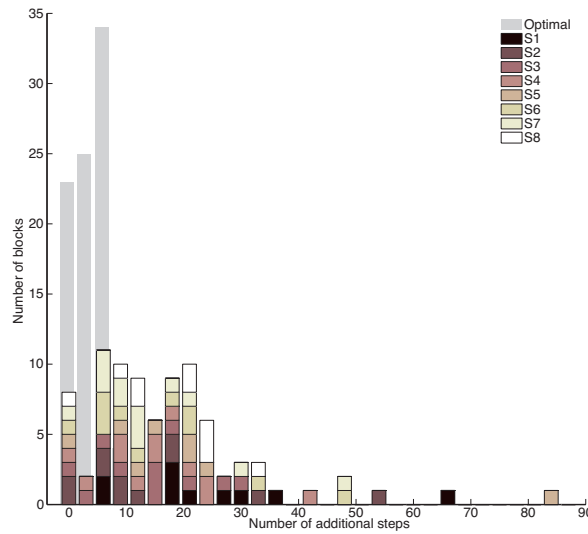


FIGURE 3.2: Histogram of additional steps needed by human participants, in comparison to the full-observation agent. The human participants took more steps to reach the goal (colored bars) than the optimal agent (gray bars).

3.2.3 Assumptions on human model

The basic assumption here is that the human participants would like to behave like as the optimal agent, but indeed behaved sub-optimally due to the limited reaction time and restriction in the memory resource. A preliminary analysis revealed that the participant's behaviors were actually different from those by the optimal agent (Fig. 3.2). Although this partially-observable maze was designed so that participants can identify their locations after at most four (meaningful) actions starting from any location on the maze, they often took non-optimal actions even after fifth trials. Then, I modeled the sub-optimal decision making by the participants, such to include the uncertain (probabilistic) character of the information processing, by means of a statistical model.

To get to the goal after as few trials as possible, identification of the present location is the key. Therefore, the participants, being familiar with the maze structure, were assumed to employ two different strategies depending on their situation; one is an exploration strategy (to try to know the present location), and the other is an exploitation strategy (to follow the shortest route to the goal). Although the same strategies have also been implemented in the optimal agent, the exploration and

exploitation by human participants are supposed to be sub-optimal, even though the available information is common with that for the optimal agent.

It would be plausible to assume that each participant keeps and maintains a belief state, sufficient statistics of unknowns, which comprises two factors: where the participant thinks he/she is (present location estimated by the participant) on the maze and whether the participant is confident of her/his estimation or not (confidence for estimation). Like the optimal agent, the participant was assumed to take an ‘info-max’ behavior if he/she is not convinced of his/her position, $c = 0$, or an optimal behavior if he/she is convinced, $c = 1$, even if those behaviors could be sub-optimal; this is the model of decision making by the participant. After getting the new observation, the participant would check if his/her previous guess was correct or not, which would make the confidence to change. Then I formulated these processes in terms of hidden Markov model (HMM).

- Internal state $H_t = (s_t, c_t)$
- Present location estimated by the participant: s_t , real location: s_t^*
- Confidence for estimation: c_t
- Actual action taken by the participant: a_t
- Actual observation: o_t^*
- Transition probability of the internal state: $p(H_t|H_{t-1}, a_t - 1, o_t^*)$
- Action selection probability: $p(a_t|H_t)$
- Initial belief $p(H_1|o_1^*)$

Figure 3.3 shows the graphical model of the HMM, which expresses the dynamics of the subjective internal states and objective states of the participants in the maze environment.

Action selection probability

Exploration strategy. $c_t = 0$ denotes that the participant is not certain about the current position on the maze. The optimal strategy in such a case is to take the ‘info-max’ (or exploratory) behavior; that is, to reduce the number of possible locations on

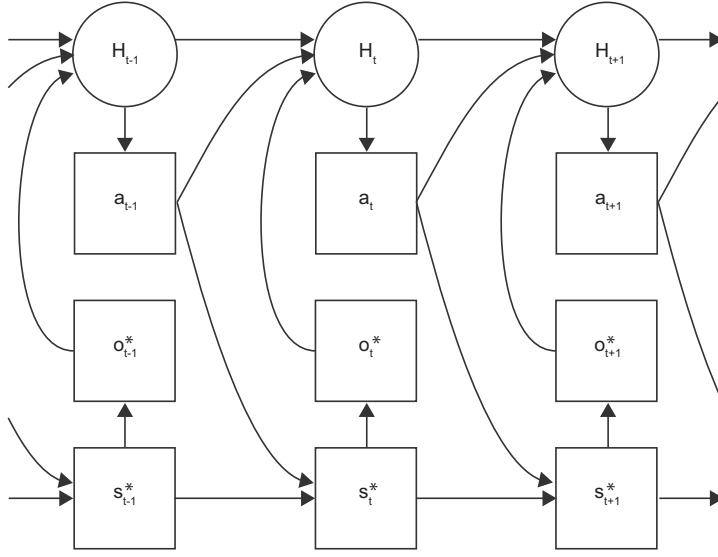


FIGURE 3.3: An HMM for human information processing during navigation behaviors. The participant's action a . is generated from the hidden internal state H . s^* denotes the real states, which involves the participant's location and orientation.

the maze. In typical situations to explore the maze, the subject is able to observe the status of three new squares by taking a forward movement, or that of two squares by taking a turn movement. Then, the exploratory behavior is assumed as follows.

1. If possible, the participant takes the forward movement, to get much new information.
2. If there is a wall in the forward direction but no wall on both of left-hand side and right-hand side, the participant takes either of turn-to-left or turn-to-right.
3. If the participant is on the L-shape corner, he/she takes the single effective turn action.

To represent the sub-optimality of the exploration by the participants, k_{exp} , which is the preference ratio of exploratory action to non-exploratory one, is introduced so that the action selection probability is given by

$$p(a_t | s_t, c_t = 0) = \begin{cases} \frac{k_{exp}}{k_{exp}N_t^{exp} + (N_t^{eff} - N_t^{exp})} & \text{if } a_t \text{ is exploratory} \\ 1 & \text{otherwise} \end{cases} \quad (3.3)$$

Here, N_{eff} and N_{exp} denote the numbers of all effective actions and exploratory actions defined above, respectively, at s_t .

Exploitation strategy. $c_t = 1$ denotes that the participant is convinced of her/his current location on the maze. In this case, the optimal behavior is to follow the shortest route to the goal. Due to the sub-optimality of the participant's exploitation, however, he/she may take some random actions with a probability ϵ_{opt} .

$$p(a_t | s_t, c_t = 1) = \begin{cases} \frac{(1 - \epsilon_{opt})}{N_t^{opt}} & \text{if } a_t \text{ is exploitative} \\ \epsilon_{opt} & \text{otherwise} \end{cases} \quad (3.4)$$

Here, N_{opt} denotes the number of optimal actions at the t -th trial. Notice that it is difficult to discriminate between the participant's exploratory and exploitative strategies only from single actions, because some exploratory actions may also be exploitative and vice versa.

Transition probability. The participants were assumed to predict the observation obtained in the next trial. Here, $\hat{s}_t$ and $\hat{o}_t$ denote the predicted position due to mental simulation (deterministic) based on the previous position estimate s_{t-1} and the actually taken action a_{t-1} , and the predicted observation at the predicted position, respectively, at the t -th trial. Due to the limited reaction time, the participant would perform mental simulation before getting a new observation at the t -th trial.

I assumed that the participant reexamined his/her position estimate based on the newly available observation; that is, if the new observation was what he/she had expected, he/she became more confident, while if the new observation was found to be different from what he/she had expected, he/she retried the estimation based on the new observation but with low confidence.

More concretely, if the participant was convinced at the $(t - 1)$ -th trial, $H_{t-1} = (s_{t-1}, c_{t-1} = 1)$, and there was no discrepancy between the predicted observation and the actual observation, $\hat{o}_t = o_t^*$, the estimated location changed to $\hat{s}_t$ with probability 1; that is, the internal state changed to $\hat{H}_t = (\hat{s}_t, c_t = 1)$. But if not convinced at the $(t - 1)$ -th trial, $H_{t-1} = (s_{t-1}, c_{t-1} = 0)$ and $\hat{o}_t = \hat{o}_t^*$, he/she remained uncertain with probability $1 - \epsilon_H$ or became convinced with ϵ_H ; that is, H_{t-1} changed to $\hat{x}_t = (\hat{s}_t, c_t = 0)$ with probability $1 - \epsilon_H$ or $(\hat{s}_t, c_t = 1)$ with ϵ_H ; the latter process represents the change of belief into a confident one.

On the other hand, if there was discrepancy between the prediction and actual, $\hat{o}_t \neq o_t^*$, the new position estimate was uniformly taken from possible positions on the maze whose view was identical to o_t^* with low confidence; that is, H_{t-1} changed to $\hat{H}_t = (s'_t, c_t = 0)$. Here, s'_t denotes any possible location whose view is identical to o_t^* .

Initial belief

At the start point, the participants were assumed to be unconvinced of their position, so $p(c_1 = 0) = 1$. Then $p(s_1)$ was a uniform distribution over the positions on the maze whose observation was identical to the initial observation o_1^* .

Estimation and prediction

Estimation for the belief state. The incremental Bayesian estimation provides a way to estimate the belief state along the series or trials (Kitagawa and Sato, 2001). Here, a belief state is the filtered estimation of the participant's internal state at the t -th trial, based on the sequences of actions and observations available for the participant before that trial. The action probability $p(a_t|H_t)$ was given in Section 3.2.3, and the state transition $p(H_t|H_{t-1}, a_{t-1}, o_t^*)$ was described in Section 3.2.3.

$$p(H_t|a_{1:t}, o_{1:t}^*) = \frac{p(a_t|H_t) \sum_{H_{t-1}} p(H_t|H_{t-1}, a_{t-1}, o_t^*) p(H_{t-1}|a_{1:t-1}, o_{1:t-1}^*)}{p(a_t|a_{1:t-1})} \quad (3.5)$$

On the other hand, by using the whole sequence of actions and observations of the participant until the end of that block, I, as an experimenter, can obtain the smoothed estimate of the belief state, which would be more reliable than the filtered estimate as follows. Here, T denotes the trial length in the block.

$$p(H_t|a_{1:T}, o_{1:T}^*) = p(H_t|a_{1:t}, o_{1:t}^*) \sum_{H_{t+1}} \frac{p(H_t|H_t, a_t, o_{t+1}^*) p(H_t|a_{1:T}, o_{1:T}^*)}{p(H_{t+1}|a_{1:t}, o_{1:t}^*)} \quad (3.6)$$

Prediction of action. The incremental Bayesian formulation also provides the way to predict the next action based on the past data as follows. The prediction accuracy is not only the measure of model's predictability (the larger, the better), but also the

quantity signifying subjective predictability of human participants.

$$\hat{a}_t = \arg \max_{a_t} p(a_t | a_{1:t-1}) \quad (3.7)$$

The concordance rate ρ between the predicted actions $\hat{a}_t$ and the real actions a_t was calculated as $\rho = N_\rho / (T - 1)$. N_ρ is the number of trials whose actions were well predicted, $\hat{a}_t = a_t$, in the block $t = 1 : T - 1$. When there were multiple probable actions (given by Equation (3.7)) as $\hat{a}_t$, I regarded $\hat{a}_t = a_t$ when a_t corresponded to any of those probable actions. In order to evaluate the HMM-based behavioral model, its prediction ability was compared with that by the optimal agent. I could evaluate the concordance rate of the optimal agent, ρ_{opt} , in a similar fashion to that of the HMM-based model.

Parameter estimation. According to the incremental Bayesian estimation, the marginal likelihood was automatically calculated, which could in turn be used for the maximum likelihood (ML) estimation of the parameter $\epsilon \equiv (k_{exp}, \epsilon_{opt}, \epsilon_H)$. Here, the parameter estimation was crucial because the parameter would represent the character of each participant. Given a set (N sequences) of actions and observations of the participant for the parameter estimation, $b_k (k = 1, \dots, N)$, the total marginal likelihood is given by

$$L(\epsilon) = \prod_{k=1}^N p(b_k | \epsilon), \quad (3.8)$$

where the marginal likelihood of a single sequence (block) $p(b | \epsilon)$ was simply given as a normalization constant of the filtered posterior of the internal state at the end of the block:

$$p(b | \epsilon) = p(a_1 | \epsilon) \prod_{t=2}^N p(a_t | a_{1:t-1} | \epsilon) \quad (3.9)$$

The ML estimate of the parameter is given by maximizing the total marginal likelihood. Although I can analytically obtain the ML estimate in this case of HMM, I simply applied the grid search heuristics: I discretized each of the parameters, k_{exp} , from 2 to 97.44 by a geometric order (common ratio 1.54, from the first to ninth terms), ϵ_{opt} , from 0.45 to 0.00 by -0.05 and ϵ_H , from 0.1 to 1 by 0.1, evaluated the total marginal likelihood on each of the $10 \times 10 \times 10$ grid points, and selected the best point that maximizes the total likelihood.

TABLE 3.2: Results of parameter estimation for the HMM (mean $\pm$ std across leave-one-block-out CV in each participant.).

Participant	k_{exp} (range 2.00-97.44)	ϵ_{opt} (range 0.45-0.00)	ϵ_H (range 0.10-1.00)
1	17.58 $\pm$ 0.00	0.05 $\pm$ 0.00	0.10 $\pm$ 0.00
2	100.00 $\pm$ 0.00	0.47 $\pm$ 0.01	0.10 $\pm$ 0.00
3	82.37 $\pm$ 18.41	0.00 $\pm$ 0.00	0.10 $\pm$ 0.00
4	55.97 $\pm$ 11.56	0.00 $\pm$ 0.00	0.10 $\pm$ 0.00
5	40.28 $\pm$ 4.93	0.10 $\pm$ 0.02	0.10 $\pm$ 0.00
6	70.62 $\pm$ 13.72	0.00 $\pm$ 0.00	0.10 $\pm$ 0.00
7	100.00 $\pm$ 0.00	0.00 $\pm$ 0.00	0.10 $\pm$ 0.00
8	52.30 $\pm$ 11.92	0.00 $\pm$ 0.00	0.10 $\pm$ 0.00

3.3 Results

3.3.1 Parameter Estimation

I performed the parameter estimation for each participant; the ML estimate of the parameter was obtained by using the data in the three sessions. The ML estimate was $\epsilon = (k_{exp}, \epsilon_{opt}, \epsilon_H) = (17.58 - 100.00, 0.00 - 0.47, 0.10)$ (Table 3.2).

3.3.2 Action prediction

Figure 3.4a shows an example of the real trajectory of a participant during a single block, starting from the red arrow (hidden start state) and ending at the double circle (instructed goal position). When the computer agents, the full-observation and optimal agent, navigated this environment, the total steps were fewer than that of the participant (Fig. 3.4b, c). The stroll agent walked around and around, and took more steps than those by the human participant. I calculated the concordance rate of all actions for each participant. The result is shown in Table 3.3, where the concordance rate ρ by the HMM was better than that by the optimal agent (ρ_{opt}) and the stroll agent (ρ_{stroll}) for almost all the participants. When there were two or three meaningful actions (decision trial, i.e., a T-junction), the action prediction was more essential than that in a straight road or a L-shape corner. Thus, I additionally calculated the concordance rate in the decision trial (Table 3.4). These results show that the HMM-based behavioral model well reproduced the participant's behaviors. When calculating the concordance rate, the parameter was estimated by using the

TABLE 3.3: Action concordance rate by the HMM (ρ), that by the optimal agent (ρ_{opt}) and the stroll agent (ρ_{stroll}).

Participant	$\rho(\%)$	$\rho_{opt}(\%)$	$\rho_{stroll}(\%)$
1	82.4	67.2	78.5
2	87.1	81.2	81.4
3	91.0	76.8	89.9
4	88.3	75.8	86.1
5	89.3	71.5	86.2
6	86.3	76.2	78.7
7	90.5	82.2	86.4
8	90.8	76.6	86.1

TABLE 3.4: Action concordance rate in the decision trial.

Participant	$\rho(\%)$	$\rho_{opt}(\%)$	$\rho_{stroll}(\%)$
1	77.1	59.2	68.2
2	80.2	77.6	69.4
3	85.5	67.5	82.2
4	82.6	68.7	77.7
5	87.9	68.0	81.5
6	79.4	72.2	66.3
7	85.4	76.8	77.2
8	87.3	72.5	78.8

data during $M - 1$ blocks out of M blocks (M is the total number of blocks, training data), and tested by using the remaining one block (test data), for each participant; i.e., the prediction performance is a cross-validated one.

3.3.3 Belief state estimation

In the early part of the block in Fig. 3.4a, the HMM estimated that the participant might have recognized his/her position but with low confidence ($c_t = 0$). As the trials proceeded, the HMM estimated that she/he was confident of his/her position ($c_t = 1$, cyan open circle), which shows good agreement with the participant's report (green circle in Fig. 3.4a). This example typically shows that the HMM well reproduced the participant's internal processing of resolving uncertainty (his/her position on the maze) in the environment, only from the participant's behaviors. In order to evaluate further the HMM-based behavioral model, I next calculated

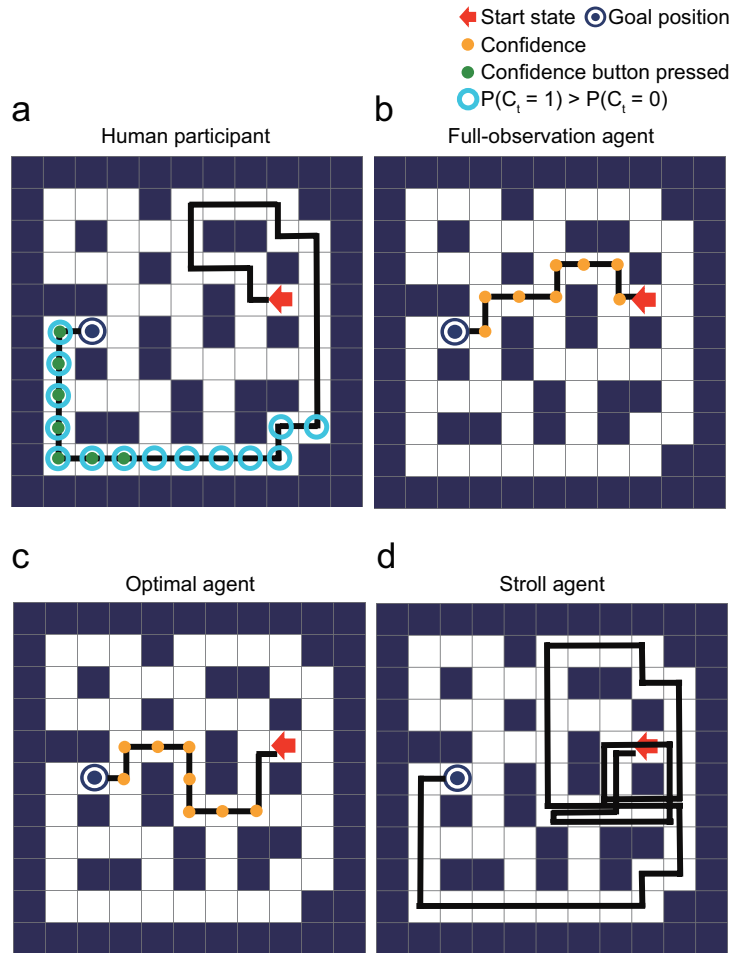


FIGURE 3.4: Navigation behaviors of a human participant (a), the full-observation agent (b), the optimal agent (c), and the stroll agent (d). Yellow circles indicate confidence of the current state. Because the full-observation agent knows the current state from the initial trial, the agent has the confidence in all trials. The optimal agent detects the current state after several steps of exploration. On the other hand, the stroll agent can not estimate the current state and has no confidence. Thus, a trajectory of the stroll agent looks like wondering along the map. Cyan open circles show trials with confidence simulated by the HMM.

TABLE 3.5: The true positive rate of the participants' confidence (ρ_{con}) by the HMM.

Participant	$\rho_{con}(\%)$
1	71.1
2	61.6
3	78.4
4	80.4
5	72.3
6	58.5
7	62.6
8	73.5

the concordance rate of confidence prediction ρ_{con} , which is the rate of the trials at which the participants reported their successful prediction to the trials at which the model evaluated that the participants were confident; i.e., true positive rate. Here, I defined the criterion of the confidence level; if the probability $\max_{s_t} p(s_t, c_t = 1)$ was bigger than $\sum_{s_t} p(s_t, c_t = 0)$, I regarded the participant as being confident of his/her estimation, otherwise as uncertain. The result in Table 3.5 shows that whether he/she was confident or not was well predicted by the HMM-based model.

3.4 Discussion

In this chapter, I presented an HMM-based behavioral model of human behaviors during performing a partially observable maze navigation task, which directly implemented switching between exploration and exploitation. The high prediction ability of the participant's actions and sufficient reproducibility of the participants' confidence suggested the plausibility of the behavioral model; that is, the participants actually performed exploration and exploitation to appropriately adapt to this uncertain environment.

I used the HMM with three parameters; k_{exp} , ϵ_{opt} and ϵ_H . I estimated these model parameter by the ML. Table 3.2 shows that these parameters were close between participants. This result implied a reliability of the parameter estimation.

Yoshida and Ishii (2006) demonstrated the brain imaging analysis based on a Bayesian HMM. My HMM described the participants' internal confidence as $P(c_t = 1)$ and $P(c_t = 0)$. One possible future study would be a model-dependent

analysis of brain activity when human participants are performing goal-directed navigation behaviors in this kind of partially observable environment.

Chapter 4

Decoding View Anticipation

Human behaviors involved in navigation are composed of two typical behavioral patterns: an exploratory action and an exploitative action (Chapter 3). When selecting an exploitative action, humans should be confident of their states in the current environment. If we call the reliable estimate of the current state of the environment as ‘belief’, they should have belief. In this chapter, I demonstrated decoding of belief from human fMRI activity patterns.

4.1 Experimental setting

4.1.1 Participants

Eight healthy participants (1 female author, 7 male; aged 21–28 years) took part in this experiment. Each participant gave written informed consent, and all experiments were approved by the Ethics Committee of the Advanced Telecommunications Research Institute International, Japan. All methods were carried out in accordance to the approved guidelines. I verified that the main decoding results remained unchanged even if the author was removed from the participant set by a post-hoc analysis.

4.1.2 fMRI Data acquisition

A 3.0-Tesla Siemens MAGNETOM Trio A Tim scanner was used to acquire interleaved T2*-weighted echo-planar images (EPI) (TR = 2 s, TE = 30 ms, flip angle = 80°, matrix 64×64, field of view 192×192, voxel 3×3×4 mm, number of slices 30). A high-resolution T1 image of the whole head was also acquired (TR = 2250 ms, TE = 3.06 ms, flip angle = 90°, field of view 256 × 256, voxel 1 × 1 × 1mm).

4.1.3 Behavioral task

I used a passive navigation game which is modified from the partially observable navigation game (Chapter 3). The experimental task (Fig. 4.1) consisted of 320 trials, spread out over five sessions (64 per session). The sessions were separated by breaks of a few minutes each. Each trial consisted of four periods. In the first period (map period), participants were presented with an initial position and orientation on one of three 2D maps (2 s). The maps had 7×7 squares of white paths and black walls with no dead ends, cross roads, white patches (2×2 squares consisting only of paths), black patches or checker patterns (Fig. 4.1). The borders of the three maps were all walls. The topographies of the three maps were symmetrical and simple so that it was easy for the participants to memorize the topographies but still hard to identify their position from only a single 3D scene. A 3D scene in wire-frame form of the wall status (path/wall) of the left, right, forward-left, forward-center and forward-right squares of the current state was presented (i.e., was partially observable). During the second period (move period), each participant's state was moved automatically by three steps (1 s/move). For each movement, after being presented with a white arrow on the current scene view giving a preview of the next move, a new 3D scene with a new preview was presented. The first and second movements were one of three movement types (move forward or turn left/right), while the third movement was fixed to be forward. Before the third movement but after presenting a preview of the third move, there was a delay period (5 s, 60 trials/session). The delay period was either 1s, 3s or 5s in each trial, to introduce 'jitter' to the raise in the brain activity; the order of this jitter time was pseudo-randomized across trials. In the following analyses, I only used the trials, 60 trials/session, whose delay time was 5s). After the delay period, the participants were then requested to choose as soon as possible between two options, a correct scene to be seen after the third movement and a distracter scene, both represented by 3D wire-frame scenes, using an MRI compatible button box (choice period, < 1.8 s). The new 3D scene after the third forward movement contained three newly observed view parts (forward-left, forward-center and forward-right), which had never been seen in the preceding move period or in the delay period. A distracter scene was one of the upcoming three view parts flipped without producing a new dead end, cross road, white/black patch or checker pattern, where flipping was balanced between the upcoming three

view parts, so the participants could not predict the scene choice task. After the participants made a choice (or if they failed to answer within 1.8 s), the options were replaced immediately by blank and the next trial started. After each session, the rate of correct answers in the session was presented to the participants. I prepared sets of 2D maps, initial states, three movements and distracters, based on diversity, so that every movement was effective (e.g., there was no wall in the forward direction after the second movement) and there was no common tendency to the participants' state after the three movements; the rate of wall/path in the scene choice task (mean $\pm$ SD) was approximately 50% for all three view parts: forward-left 0.49 ± 0.01 , forward-center 0.50 ± 0.01 and forward-right 0.50 ± 0.02 . The SD was over five sessions. According to these criteria, 167 sets were prepared and used across the experiments. These sets were common for all the participants, but their order was randomized within a session such to be counterbalanced across participants. Here, I prioritized the unpredictability of the scene choice task, while the scene occurrence was not uniform, reflecting the map topography; some sets were presented more than once (max 5, mean 1.95).

On the day before scanning, participants were given verbal and written explanations of the aim and procedures of the experiment, and practiced two types of training sessions outside the MRI scanner. In the first training session, the participants performed a free navigation task in another map with larger 9×9 squares, in which they selected their movements by their own, so as to learn the relationship between 2D topography and 3D scenes. Participants were presented with both the current state on a 2D map and the 3D scene of the current state. After the participants reported that they were familiar with the 2D-3D association, in the second type of training they performed sessions of a scene choice task, which was the same as the main experimental task that they would take part in the subsequent day. This second type of training was ended after the participants could choose the true upcoming scene with more than 80% accuracy.

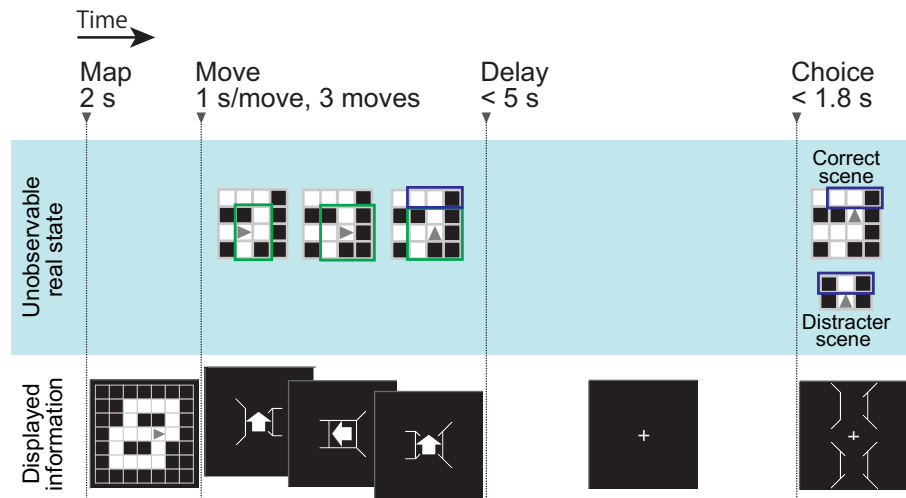


FIGURE 4.1: An example scene choice trial. The bottom display series show an example of presented stimuli during a single scene choice trial. In the first display, one of the three 2D maps (see Subsection 4.1.3) and a triangle indicating the initial state, position and body orientation, of the participant on the 2D map were presented (map period). The participants' state was then moved automatically by three steps; each movement was shown by an arrow placed at the center of a 3D scene of the current state (move period). Each 3D scene showed the path/wall status of left, right, forward-left, forward-center and forward right squares of the current state (a triangle in the top 'real state' display). An up arrow made the participant to move forward with the same orientation, while a left (right) arrow made the participant's orientation to turn toward the left (right) on the same square. The second display showed the 3D scene of the initial state and first movement of the participant; in this case, the participant was moved forward to the east. The participants' orientation was changed to the north (turn left, second movement) in the third display, while the participant was moved forward (third movement) in the fourth display. The right display series show the real states of the participant, although they were hidden for the participant. Before the third movement but after presenting a preview of the third move, there was a delay period during which only a fixation point was presented (the left fifth display), then the participant was requested to choose the upcoming scene from two options: a correct scene to be seen after the movement and a distracter scene (the sixth display) (choice period). Because the third movement was fixed to be forward (the fourth display), the prediction of the upcoming scene after the third movement always constituted three new view parts, whose direct information was not seen during the move period. The distracter scene was one of the three view parts flipped (the forward-left view was turned from a path to a wall in this example).

4.2 Analysis methods

4.2.1 Preprocessing for fMRI data

The first six scans of each run were discarded so as not to be affected by initial field inhomogeneity. The acquired fMRI data underwent 3D motion correction using SPM8 (<http://www.fil.ion.ucl.ac.uk/spm>). The data were coregistered to the whole-head high-resolution anatomical image, and then spatially smoothed with a Gaussian kernel filter (FWHM, 8 mm). In a post-hoc analysis, I confirmed that a smaller (3 mm) setting of the spatial smoothing did not change any of my conclusions.

4.2.2 Regions of interest

During spatial navigation of rats, firing patterns of hippocampal system were reported to represent spatial status (Hafting et al., 2005; Harris et al., 2003). Neurons in medial prefrontal cortex (mPFC) and posterior parietal regions showed choice- and proceeding- specific firing patterns (Fujisawa et al., 2008; Harvey, Coen, and Tank, 2012). Early visual areas in the occipital cortex (OC) retain specific information about contents of visual working memory when no physical visual cue is present (Tong et al., 2012). Moreover, spatial working memory is thought to be represented in dPFC (Courtney et al., 1998). Based on these previous studies related to the view expectation, I examined six bilateral regions of interest (ROIs) in the decoding analysis: the mPFC, dPFC, precuneus, superior parietal cortex (sPC), hippocampal-parahippocampal cortex (HC-paraHC) and OC (Table 4.1). I integrated the HC and paraHC into a single region, making the size of the combined region (HC-paraHC) comparable to those of the other ROIs. Using ITK-SNAP (<http://www.itksnap.org>) (Yushkevich et al., 2006), the six anatomical ROIs were identified on the high-resolution T1 image of each participant in reference to the Automated Anatomical Labeling (Tzourio-Mazoyer et al., 2002), whose mean voxel numbers and their variances over the eight participants are summarized in Table 4.1. The fMRI signals from these ROIs then underwent quadratic polynomial trend removal and noise reduction by means of singular value decomposition ($K = 3$), and were then normalized within each session.

TABLE 4.1: Datasets were created by fMRI voxels from six brain regions defined anatomically and segmented manually. Mean number of voxels is over eight participants, and SD is the standard deviation of the voxel numbers over the participants

Region	Mean number of voxels	SD
mPFC	577.1	133.2
dPFC	470.6	127.2
Precuneus	644.8	40.6
sPC	671.4	179.8
HC and paraHC	366.5	113.1
OC	858.3	117.6

4.2.3 Labeling of fMRI data

Trial-based fMRI signals taken in the scene choice task were labeled as positive or negative depending on four types of status in the scene choice task: an upcoming scene view, an observed scene view in the third move, a position in the third move, and a map. The label for the upcoming scene view was set based on the status, path (positive, or 1) or wall (negative, or 0), associated with the three view parts of the scene choice task; i.e., forward-left, forward-center and forward-right (Fig. 4.2). Although the participants were assumed to predict the next scene but could not predict the distracter, the delay period fMRI signals involved activity related to the prediction of the upcoming three view parts. To incorporate the view part dependency obtained by the behavioral analysis (Fig. 4.3) and to avoid overly complicated procedures in the scanning experiment, the distracters were set as one-view-part flipped ones; this setting was appropriate for the construction of binary classifiers (decoders). Similarly, the label for the observed scene views consisted of two codes; whether the left and right sides at the state after the third movement were path (positive, or 1) or wall (negative, or 0), which are observable as the forward-left and forward-right view parts before the third movement. For the position, the label consisted of four binary codes, as follows: whether the position after the third movement was located in the upper side (positive, or 1) or not (negative, or 0), in the lower side or not, in the left side or not, and in the right side or not. Because each of the three maps had 7×7 squares, the codes could be positive both for the upper (left) side and lower (right) side when the position was at the middle row (column)

of the map. The map label consisted of three binary codes, each of which signified each of the three maps 1, 2 and 3; e.g., (1, 0, 0) indicates map 1.

4.2.4 Decoding analysis overview

Multiple parallel decoding analysis was performed on single trials based on the fMRI signals of the voxels belonging to the six ROIs above. Each binary decoder was trained in a supervised fashion to associate the normalized voxel signals (input variables) with the above-mentioned labels (output variables) in the scene choice task. For each binary classifier, I used a linear classifier, sparse logistic regression (SLR) (Yamashita et al., 2008). Although the number of inputs was as large as 200–1000 (Table 4.1), hierarchical Bayesian setting called automatic relevance determination and its variational Bayesian approximation can ignore input variables irrelevant to the classification; such sparseness would be effective in increasing the generalization capability. In my implementation, default values in the SLR toolbox were used for all parameters of the classifiers.

Comparison between status-specific decoders

For each participant, 12 binary SLR classifiers (three decoders for the upcoming scene view, two for the observed scene view, four for the position and three for the map), termed status-specific decoders in the multiple parallel decoding, were trained individually (Fig. 4.2). 12 binary status-specific decoders were trained for each scan and for each of the six ROIs over the single trials, and the total 72 (12×6) decoders were evaluated. When evaluating the trained status-specific decoders, I used a leave-one-session-out (LOSO) cross-validation (CV). In the LOSO procedure, each decoder was trained using the training data set from which one session was removed, and the removed session was used to validate the trained decoder. By changing the removed session one by one, I could evaluate the CV performance of the decoders. When reporting the overall decoding accuracy of each status-specific decoder (Fig. 4.4a), the mean accuracy averaged over the constituent code-dependent decoders (e.g., forward-left, forward-center and forward-right, for the upcoming-scene-view-dependent decoders) was used.

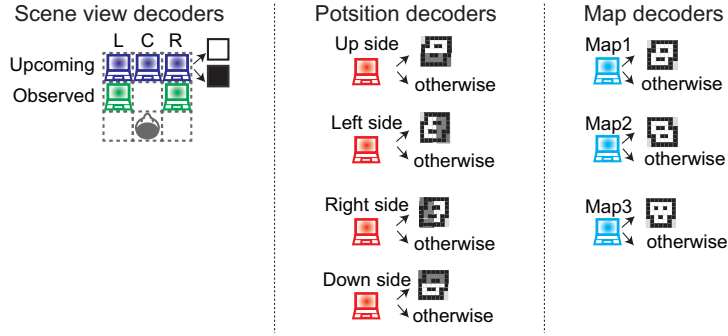


FIGURE 4.2: Analysis overview. Four types of navigation-related decoder were applied to the brain activity during the scene choice task. The scene view decoders were of three binary classifiers in terms of upcoming square: forward-left, forward-center and forward-right, and two binary classifiers in terms of observed square: left and right. The position and map decoders were for the rough position and the map, respectively, where the participant stayed in the current trial. In this example's case, the up-side, right-side and map 1 classifiers' output should be positive because the current location was upper right on map 1.

For each decoder, samples with the positive and negative labels in the original training data set were unbalanced (The ratio (and the SD over eight participants) of positively-labeled samples in the set of correct trials was as follows: the upcoming scene view, forward-left 0.28 (0.01), forward-center 0.67 (0.02) and forward-right 0.28 (0.01); the observed scene view, left 0.35 (0.02) and right 0.36 (0.01); the position, upper side 0.81 (0.03), lower side 0.68 (0.01), left side 0.59 (0.01) and right side 0.63 (0.01); and the map, map1 0.28 (0.01), map2 0.31 (0.02) and map3 0.41 (0.03)). Because simple use of such an unbalanced data set to train binary decoders may introduce a bias to the decoders, positive/negative samples in each training data set were resampled to remove the unbalance. Let N_{pos} and N_{neg} be the numbers of positive and negative samples in the original training data set, respectively. If $N_{neg} > N_{pos}$, then N_{pos} negative samples were randomly subsampled, otherwise N_{neg} positive samples were subsampled, so that the numbers of positive and negative samples became equal. An output from each SLR decoder was an analogue value, ranging between 0 and 1, implying the probability of the label being one (positive). When evaluating the decoder's accuracy, I binarized the analogue value by applying a threshold value of 0.

Evaluation of the decoders for the upcoming scene view

When evaluating each of the three decoders for the upcoming scene view, the fMRI data samples for each participant were divided into three groups in terms of the flipped view parts in the scene choice task. For each participant and for each of the six ROIs, three binary SLR classifiers were trained individually. The set of voxel signals averaged over two scans ($2 \times \text{TR} = 4 \text{ s}$) just before presenting the choice options (the choice period) was used as a sample for the decoders. The three decoders were then constructed based on the three corresponding separated data samples. Thus, the label unbalance in the training data set was minimal (prior chance levels of the decoder after removing miss trials, forward-left: $51.9 \pm 2.2\%$, forward-center: $50.9 \pm 0.5\%$, forward-right: $51.4 \pm 1.3\%$). When evaluating the three upcoming scene decoders, I used a simple leave-one-trial-out (LOTO) procedure, in which each decoder was trained by the training data set from which one trial was removed, and the removed trial was used for validation (Fig. 4.4b, 4.4c).

4.2.5 Reconstruction methods

When visualizing the maps, I used an alternative validation method termed the leave-one-map-out (LOMO) procedure. Each decoder for a single view part (either of forward-left, forward-center or forward-right) was trained using the data set in the scene choice task of the corresponding view part, but the training trials were restricted to those using two of the three maps. The trials employing the remaining map were kept for validation. The objective of the LOMO procedure was to examine the generalization capability of the decoder's neural basis by removing the dependency on specific maps. In the visualization, I calculated the square-wise value (corresponding to the probability of path of that square), $P_{X,Y,M}$, as the mean

of the decoders' analogue outputs $\hat{C}$ with the sigmoidal rule:

$$\begin{aligned}
 C_{f\star,i} &= \frac{\exp(\beta \hat{C}_{\star,i})}{\exp(\beta \hat{C}_{\star,i}) + \exp(\beta(1 - \hat{C}_{\star,i}))} \\
 P_{X,Y,M} &= \frac{\sum_{i=1}^N \delta_{fL,i} C_{fL,i} + \sum_{i=1}^N \delta_{fC,i} C_{fC,i} + \sum_{i=1}^N \delta_{fR,i} C_{fR,i}}{\sum_{i=1}^N \delta_{fL,i} + \sum_{i=1}^N \delta_{fC,i} + \sum_{i=1}^N \delta_{fR,i}} \\
 \delta_{f\star,i} &= \begin{cases} 1 & x_{f\star,i} = X, y_{f\star,i} = Y, m_{f\star,i} = M \\ 0 & \text{otherwise} \end{cases}
 \end{aligned}$$

where β is an inverse-temperature parameter ($\beta = 5$), $x_{f\star,i}$ and $y_{f\star,i}$ represent the horizontal and vertical coordinates, respectively, of the target square of the decoder either of forward-left, forward-center or forward-right in the i -th trial and m_i is the map index of the trial i . N is the number of trials in which the map M was used by the participant. The averaged numbers of the superposition ($\sum_{i=1}^N \delta_{fL,i} + \sum_{i=1}^N \delta_{fC,i} + \sum_{i=1}^N \delta_{fR,i}$) were 6.1 ± 6.8 for the squares on map 1, 6.5 ± 5.9 for those on map 2 and 8.7 ± 4.0 for those on map 3 (the SD is over squares). I applied a map-dependent threshold (see legend of Fig. 4.5b) to the analogue value above to make the visualized map a binary one.

4.3 Results

4.3.1 Behavioral results

Eight participants performed the scene choice task that was a modified version of the pre-learned maze navigation game presented previously (Yoshida and Ishii, 2006); the participants were requested to choose as soon as possible between two options represented by 3D wire-frame scenes, a correct upcoming scene and a distracter scene (Fig. 4.1). The options appeared after the participants' state was moved automatically in one of three two-dimensional (2D) pre-learned mazes by three steps and before seeing the scene after the third movement. The scene choice task required the participants to identify small differences between the two options with only one of three view parts, forward-left (L), forward-center (C) or forward-right (R) flipped (Figs. 4.1 and 3, subsection 4.1.3). Although an allotted time was limited (choice period, 1.8 s) and any information about the scene choice task was not presented until

the choice period, the participants showed satisfactorily high performance (mean $\pm$ SD): $93.4 \pm 2.6\%$ correct, $5.6 \pm 2.3\%$ incorrect and $1.0 \pm 1.0\%$ missed. Missed trials, in which the participants could not press an answering button in the allotted time, were excluded from the following analyses. To answer correctly and rapidly, the participants were required to predict the upcoming scene, and the behavioral result above showed that participants predicted the upcoming scene well. A Friedman test for participants' task accuracy and reaction time (RT) demonstrated significant differences between flipped view parts (Fig. 4.3a; task accuracy, $\chi^2(2, N = 40) = 33.7, p < 0.05$; RT, $\chi^2(2, N = 40) = 9.2, p < 0.05$); in the forward-center flipped trials, the participants had a tendency toward higher accuracy and shorter RT than in other trials. By contrast, there were no apparently significant differences in the task accuracy and RT between the three maps (Fig. 4b; task accuracy $\chi^2(2, N = 40) = 5.9, p = 0.05$, $\chi^2(2, N = 40) = 0.1, p = 0.95$). Additionally, the task accuracy did not show significant differences between the rough position on the maze (Fig. 4.3c; Wilcoxon signed rank test, up-side vs. otherwise: $p = 0.74$; down-side vs. otherwise: $p = 0.95$; left-side vs. otherwise: $p = 0.08$; right-side vs. otherwise: $p = 0.25$), between the statuses (wall/path) of the current observable scene (Fig. 4d; Wilcoxon signed rank test, left: $p = 0.74$, right: $p = 0.38$) or between the positions (up/down, see Fig. 4.2) of the target option (correct scene) (Fig. 4.3e; Wilcoxon signed rank test, up-side vs. down-side: $p = 0.30$). In addition, there was no significant difference in the RT's distribution between correct trials and incorrect trials (Fig. 4.3f; Wilcoxon rank sum test, $p = 0.45$); although the RT would reflect the difference in cognitive processes between view parts, it has no information about which option was chosen by the participants in each trial.

4.3.2 Decoding performance

Time-course decoding analysis

To determine when and in what region the neural bases were recruited in the scene choice trials, I performed voxel-wise fMRI analysis along the whole task. I constructed four sets of binary decoders (Fig. 4.2, Subsection 4.2.4): three upcoming scene view decoders, two decoders for observed scene view in the third move, four decoders of position in the third move, and three map decoders. Each decoder associated voxel-wise activity patterns in each of six anatomical regions of interests

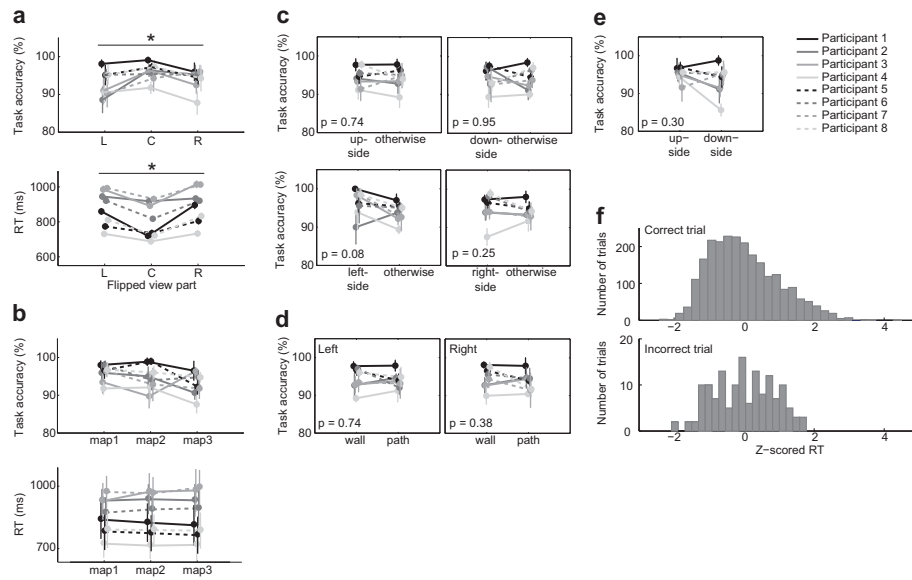


FIGURE 4.3: Behavioral results. They are shown in terms of task accuracy (how accurately the participants performed the scene choice task) and reaction time (RT), both averaged over sessions. Error bars indicate SEM across sessions ($n = 5$). (a) In the forward-center (C) trials, the task accuracy was higher and the RT was shorter than in the forward-left (L) and forward-right (R) trials (Friedman-test, $p < 0.05$). (b) No significant difference in the task accuracy or RT was found between three maps (Friedman-test). (c) The position of the participants after the third movement did not significantly affect the task accuracy (Wilcoxon signed rank test). (d) The task accuracy did not show significant difference between the statuses (wall/path) of both of the left and right sides of the observed scene view (Wilcoxon signed rank test). (e) The position, up or down, showing the correct scene in the choice period did not significantly affect the task accuracy (Wilcoxon signed rank test). (f) Histograms of RTs in correct trials (top panel) and in incorrect trials (bottom panel). Each RT was converted to z score per session.

(ROIs): mPFC, dPFC, precuneus, superior parietal cortex (sPC), HC-paraHC and OC (Table 4.1, Subsection 4.2.2). fMRI data preprocessing included trend removal, noise reduction and temporal normalization (see subsection 4.2.1). Here, any hemodynamic response function was not convolved to avoid confounding any prospective activation related to the behaviors (option choices) induced by visual presentation in the subsequent choice period. When I used trials in which the participants answered correctly in the scene choice task (correct trials) for both of decoders' training and their evaluation, I found that decoding of the upcoming scene view was possible from the single-scan (2 s) activities in the mPFC, precuneus, sPC and OC (Fig. 4.4a). Interestingly, the decoding accuracy from those regions was high especially in the delay period, and was even better in the choice period. Activities in the sPC allowed me to decode the observed scene view in addition to the upcoming scene view: for the fourth scan (move period), 53.2%, $p < 0.05$ (with Bonferroni correction for comparing multiple time points and decoders); for the fifth scan (early delay period), 53.7%, $p < 0.05$ (Bonferroni-corrected). Most notably, activities in OC allowed us to decode the position in addition to the upcoming scene view: for the second scan (map period), 56.2%, $p < 0.05$ (Bonferroni-corrected); for the third scan (move period) 58.2%, $p < 0.05$ (Bonferroni-corrected); for the fourth scan (move period), 54.3%, $p < 0.05$ (Bonferroni-corrected); for the eighth scan (choice period), 52.7%, $p < 0.05$ (Bonferroni-corrected). However, there were no significantly decodable scan timings when using dPFC or HC-paraHC activities. These decoding results remained unchanged even if I used the activities in HC or paraHC alone.

Upcoming scene view decoding from delay-period brain activity

To study responses of brain activities to the view expectation in terms of view parts, I evaluated the three upcoming-view-part-dependent decoders with correct trials. In each trial of the scene choice task, the participant's choice from two options allowed me to know what had been the participant's expected scene for a single view part (corresponding to the flipped view part in the distracter). To construct effective decoders that exploited this character, I divided the fMRI data samples into three groups in terms of the flipped view parts, so that three view-part-dependent decoders (forward-left, forward-center and forward-right) were individually trained based on the three groups. Their performance was examined by the leave-one-trial-out (LOTO) procedure (Subsection 4.2.4), with Bonferroni correction to deal with

multiple anatomical ROIs. Based on the time course analysis above, I used the fifth scan (early delay period) and the sixth scan (just before the choice period); i.e., when the participants were preparing to see the options to be chosen (Fig. 4.1).

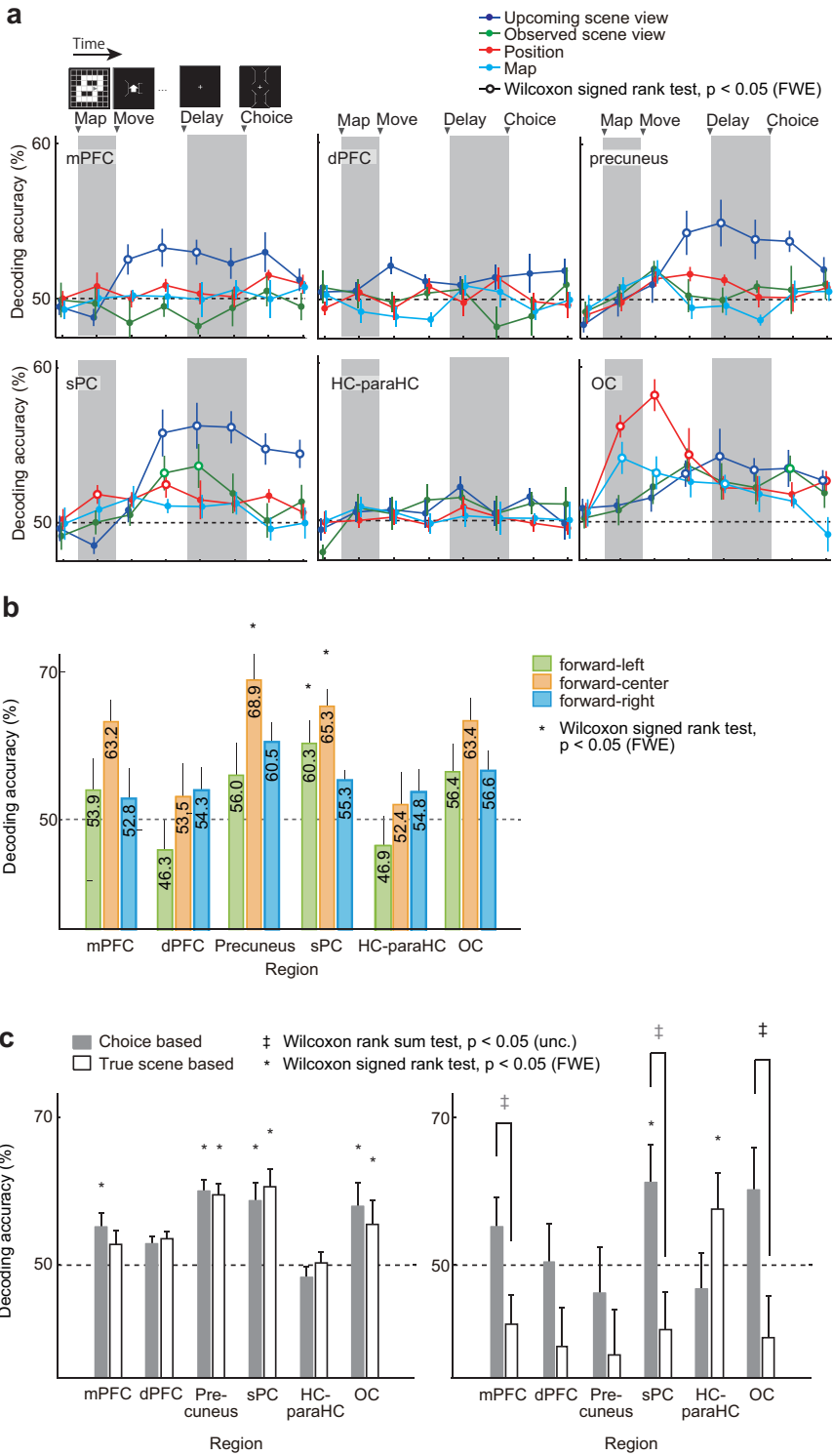


FIGURE 4.4: Decoding results. Error bars indicate SEM across participants ($n = 8$). (a) Time-course decoding analysis along the scene choice task. Each panel shows the decoding accuracy from one of the six ROIs, which was evaluated by the leave-one-session-out (LOSO) procedure. At each event timing, fMRI signals during a single scan (2 s) just after that timing in the time-course were used for decoding. Open circle indicates significantly higher accuracy than the chance level. (b) The decoding accuracy was evaluated when the trials were separated into three categories (forward-left, forward-center and forward-right), depending on the target (the pair of the correct view and the distracter view) in the scene choice task; the decoding accuracy was evaluated by the LOTO procedure. (c) The decoding accuracy in correct (left) and incorrect (right) trials was evaluated by the decoders that were trained with incorrect trials. I used two methods: one involved decoding the distracter's view as right as the participants actually chose it, while the other involved decoding the true view as right as it should be ideally chosen. The decoding accuracy was evaluated by the LOTO procedure. Double daggers indicate significant difference in the accuracy between these labeling methods (p values evaluated by the Wilcoxon rank sum test).

4.3.3 Map reconstruction

To remove the decoder's dependency on map topography, I performed another kind of CV, leave-one-map-out (LOMO), in which each upcoming-view-part-based decoder did not use the trials when the participants were on a specific map when training, but was evaluated by those trials. Thus, each decoder could not use the topographic information of the map. This LOMO CV also allowed me to visualize the map by simply arranging the scene view predicted by the upcoming-view-part-based decoders based on the participants' fMRI signals just before seeing that view part, as a 2D representation of the prediction of the map (Fig. 4.5b, 2D arrangement of decoded wall-status). Because most squares in the visualized map show the average of outputs from multiple view-part-based decoders and multiple trials, the square-wise decoding accuracy (map visualization accuracy) of the whole map (mean map visualization accuracy: map 1: $72.7 \pm 9.2\%$; map 2: $72.8 \pm 7.3\%$; map 3: $70.0 \pm 5.7\%$; SD is over eight participants) was higher than those of view-part-dependent decoders (Fig. 4.4b).

4.3.4 Reconstruction and human behavior

The most far-reaching question I posed was whether this decoded view expectation would correlate with the participants' choice behavior. As shown in Fig. 4.5c, I found that the square-wise decoding accuracy was significantly correlated with individuals' performance of the scene choice task.

4.4 Discussion

4.4.1 Status-specific decoding analyses

My decoding analyses found that the parietal regions (precuneus and sPC) represented the prediction of upcoming scene view in the scene choice task (Fig. 4.4), implying that they are involved in navigation in partially observable navigation environments. I found that various kinds of navigation-related information (i.e., the observed scene view, the map and the position), as well as the upcoming scene view, were decoded from the OC in the delay period, suggesting that the visual system in the OC is used for prediction accompanied by collecting information from observations. These results are consistent with previous decoding studies on visual working memory, mental rotation and perceptual speed (Albers et al., 2013; Ester, Serences, and Awh, 2009; Harrison and Tong, 2009; Vintch and Gardner, 2014). Although decoding performance was poor from the hippocampal system (HC-paraHC) and dorsal part of PFC (dPFC) throughout the task, the mPFC and precuneus showed good decoding performance only for the upcoming scene view. Notably, the sPC showed the highest decoding accuracy for the upcoming scene view (Fig. 4.4a); when I decomposed the scene view to be predicted into three view parts (forward-left, forward-center and forward-right), decoding from the sPC showed significantly high accuracy. These results are consistent with the findings that sPC contributes to the manipulation of egocentric spatial information (Vallar et al., 1999; Galati et al., 2000).

The time-course decoding analysis (Fig. 4.4a) provided information on the processes during the scene choice task; in the move period, prior to precuneus and sPC, the upcoming scene view was represented in the mPFC. According to (Yoshida and Ishii, 2006), in the mPFC, multiple proposals for the current position are maintained and resolved as the information from the observations increases. Thus, brain

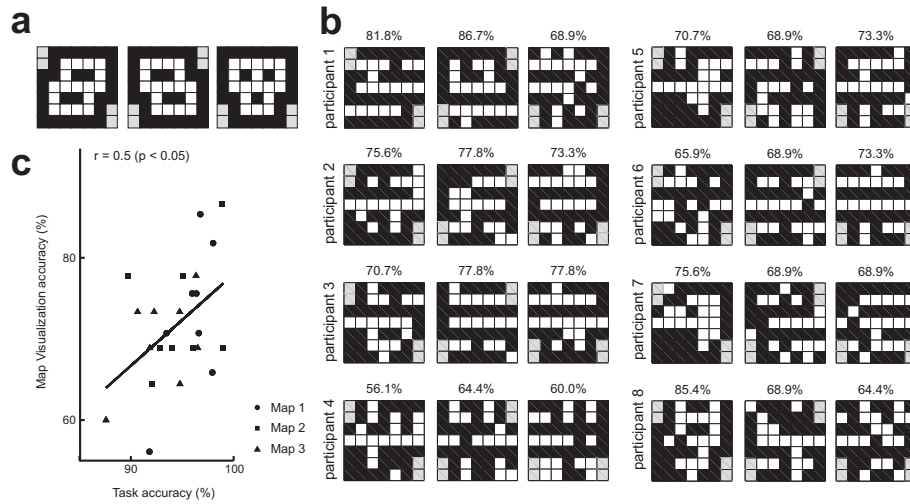


FIGURE 4.5: Real maps and visualized maps based on decoding of view expectation from fMRI activity patterns. (a) Three 7×7 maps used in the experiment. (b) Each square of the visualized map denotes the average of the outputs from the three upcoming-scene-view-dependent decoders (forward-left, forward-center and forward-right) as white (black) if the averaged value was larger (smaller) than a threshold, although the number of samples in which the corresponding square was the target of prediction varied depending on the squares. When visualizing a specific map, fMRI data and the true labels (path/wall status) in the trials where the map was used in the scene choice task were never used (the LOMO procedure), and an output from each decoder was an analogue value ranging between 0 and 1, showing the path probability (0, less path-like; 1, more path-like). Because this subjective probability can be biased because of the number unbalance between paths and walls, the threshold to binarize was set to select the same number of paths as in the real map. Gray squares were excluded from the evaluation. At the top of each visualized map, the visualization accuracy for the squares in the map is displayed. (c) Scatter plot depicts a significant relationship between the behavioral task accuracy of individual participants and their map visualization accuracy. For each participant, there are three points corresponding to the three maps.

activity in the mPFC may have contributed to retention and updating of multiple proposals, which should be preceded to the prediction of the new scene. Besides, the mPFC is considered the key nodes of the default mode network (Fox et al., 2005). The information of the upcoming scene view expectation, which was involved in mPFC activities, may thus be used for manipulating multiple proposals within the dynamics embedded in spontaneous brain activities. The position decoders showed the highest accuracy when decoded from the OC activity in the early period of the scene choice task. Because the position identification was crucial to perform well in the scene choice task on these rather symmetrical maps, a large cognitive resource might have been allotted to position identification, especially in the move period.

Because there is little known regarding the neural bases of prediction, I introduced a decomposed decoding technique, producing three binary decoders (forward-left, forward-center and forward-right), rather than a single eight ($= 2^3$)-class decoder, for the upcoming scene view. A similar method used to read out complex viewed stimuli from primary visual fields was previously reported (Miyawaki et al., 2008). Such parallel decoding methodology shows conceptual correspondence to the encoding process by the brain, which is implemented within massively parallelized machinery.

Previous decoding studies often used localization tasks to define the functional ROIs of individual participants (Harrison and Tong, 2009; Kamitani and Tong, 2005; Naselaris et al., 2009). For higher-order cognitive functions such as view expectation, however, functional ROIs are too dependent on the localization task itself, and may not be very suitable; as higher-order cognitive functions may involve multiple sub-processes, and as fMRI signals from higher-order cortices are relatively weak, it can be difficult for the localization task to identify responsible regions that are recruited only in some but essential parts of the whole task. As such, I preferred anatomical ROIs that were parcellated based only on individuals' anatomical information, rather than functional ROIs. For this purpose, I used a sophisticated software for anatomical parcellation, which was also used in a study of memory decoding from the HC and paraHC (Bonnici et al., 2012).

In the present study, I used structured maps rather than unstructured counterparts. The use of structured maps is advantageous in that they are more natural than unstructured maps, so they have a straight link toward natural navigation in

real environments that should be highly structured. Note here that my view-part-dependent decoders did not use any allocentric structure information (Comparison between status-specific decoders), although the structures in the maps could provide important clues for the participants to perform not only localization but also prediction. My time-course decoding analysis (Fig. 4.4a) actually showed the participants' orchestrated processing between the view expectation in the mPFC, precuneus and sPC, and the manipulation of visual and hence structural cues in the OC.

4.4.2 Decoding of expectation in decision making

Especially for decision-making in complicated environments, as for navigation in partially observable environments, the estimation of the current state and the prediction of the next state are essential sub-processes. Indeed, combining state estimation in terms of the 'belief state' (i.e., the posterior distribution of the current state) by means of optimal Bayesian observer and value evaluation in the space of belief states is known to produce optimal decision-making, even in partially observable environments (Kaelbling, Littman, and Cassandra, 1998). Although I focused on the prediction of the upcoming scene rather than value evaluation, the next step would be to examine the crosstalk between the scene prediction and the prediction-based decision-making. Consistent with recent findings that functional connectivity within a fronto-parietal network is correlated with the performance of primary category cueing (Bollinger et al., 2010), I found that contents in view expectation were decoded from and hence represented by fMRI activities in the mPFC, precuneus and sPC.

There are several previous reports of fMRI-voxel-based decoding in 3D navigation environments. (Hassabis et al., 2009) decoded the position of the participants from fMRI activities in the HC and paraHC when they navigated to one of four corners in one of two 3D environments by identifying their positions based on the cues placed on the wall. Related to this study, Rodriguez demonstrated the decoding of working memory of goal direction when participants navigated to a goal position (E, W or N) based on cues placed at a specific position (S) (Rodriguez, 2010); fairly high decoding performance was obtained by using voxels in the medial prefrontal gyrus (51.4%), hippocampus (47.7%) and inferior parietal cortex (49.2%). More recently, it has been reported that the location and direction during navigation are

encoded to activities in medial parietal regions (Marchette et al., 2014; Vass and Epstein, 2013). In this study, on the other hand, I found no evidence to suggest that position/map information is represented in neither of the HC-paraHC, sole HC, or sole paraHC. Because the computational cost to manipulate multiple proposals and predict a plausible next scene while resolving the proposals (these are speculated to be performed in the fronto-parietal network) was high in my task, the HC and paraHC could have shown relatively low decodable activities.

Here, I mainly focused on the fronto-parietal network, which has been considered as the source of anticipation; therefore, I excluded other brain regions that could be involved in spatial information processing: the response of dorsal premotor cortex is related to spatial attention and memory (Simon et al., 2002); the nucleus reuniens is at a passing point for spatial anticipation from the mPFC to the hippocampal system (Ito et al., 2015). Functional dissociation of view expectation from other spatial information processing remains as one key issue to be clarified. Actually, my decoders presented sustained decodability from the delay period activity of the precuneus and sPC but not from that of the dPFC (Fig. 4.4a). This is in contrast to the previous work in which, the dPFC showed sustained activity during spatial information processing involved in working memory (Courtney et al., 1998).

Chapter 5

Encoding View Anticipation

In the previous chapter, I described neural decoding of scene view anticipation. My proposed decoding model characterized neural representation for scene view anticipation, which allowed us to detect subjective errors of the participants. Here, I tried to unravel neural encoding in decision making in similar partially observable navigation environments, using a novel data-driven encoder modeling.

5.1 Experimental setting

5.1.1 Participants

Eight healthy participants (1 female author, 7 males; aged 21-28 years) took part in this experiment. Each participant gave written informed consent, and all experiments were approved by the Ethics Committee of the Advanced Telecommunications Research Institute International, Japan. All methods were carried out in accordance with the approved guidelines.

5.1.2 fMRI Data acquisition

A 3.0-Tesla Siemens MAGNETOM Trio A Tim scanner was used to acquire interleaved T2*-weighted echo-planar images (EPI) (TR = 2 s, TE = 30 ms, flip angle = 80° , matrix 64×64 , field of view 192×192 , voxel $3 \times 3 \times 4$ mm, number of slices 30). A high-resolution T1 image of the whole head was also acquired (TR = 2250 ms, TE = 3.06 ms, flip angle = 90° , field of view 256×256 , voxel $1 \times 1 \times 1$ mm).

5.1.3 Behavioral task

The participants performed two types of three-dimensional (3D) spatial navigation task on the same day, one of which, the scene choice (SC) task, has been described previously (Chapter 4) and the other, the motion decision (MD) task, is new to this study.

Pretraining

On the day before scanning, the participants performed practice tasks to memorize the map topography and to become familiar with the relationship between the two-dimensional (2D) maps and the 3D views of the maps (Fig. 5.1a), according to the procedure described previously (Section 4.1). In the first practice task, the participants performed a free navigation task in a map with 9×9 squares of white paths and black walls, which was the same map as would be used in the MD task on the subsequent day. After observing the current state on the 2D map and the 3D wire-frame view seen at the current state simultaneously, they indicated their chosen movement using one of three keys on a computer keyboard, for 'go forward', 'turn left', or 'turn right'. This practice task continued until the participants reported that they were familiar with the 2D-3D association. Then participants performed a practice SC task with three 7×7 maps, which were the same maps as would be used in the SC task on the subsequent day. This second type of practice continued until the participants could choose the correct upcoming scene with more than 80% accuracy.

Scene choice (SC) task

The following day, participants performed the experimental tasks in the scanner (Fig. 5.1b). The SC task consisted of 320 trials, spread out over five sessions (64 trials per session). The sessions were separated by breaks of a few min each. Between the third and fourth sessions, another experimental task, the MD task, was performed (Fig. 5.1c). An SC trial consisted of four periods: a map period (2 s), a move period (three moves, 1 s/move), a delay period (5 s) and a choice period (1.8 s). In the map period, participants were presented with an initial position and orientation on one of three 2D maps (Fig. 5.1b). During the move period, the participant's state was moved three steps by the computer. Hence, they passively viewed a sequence of

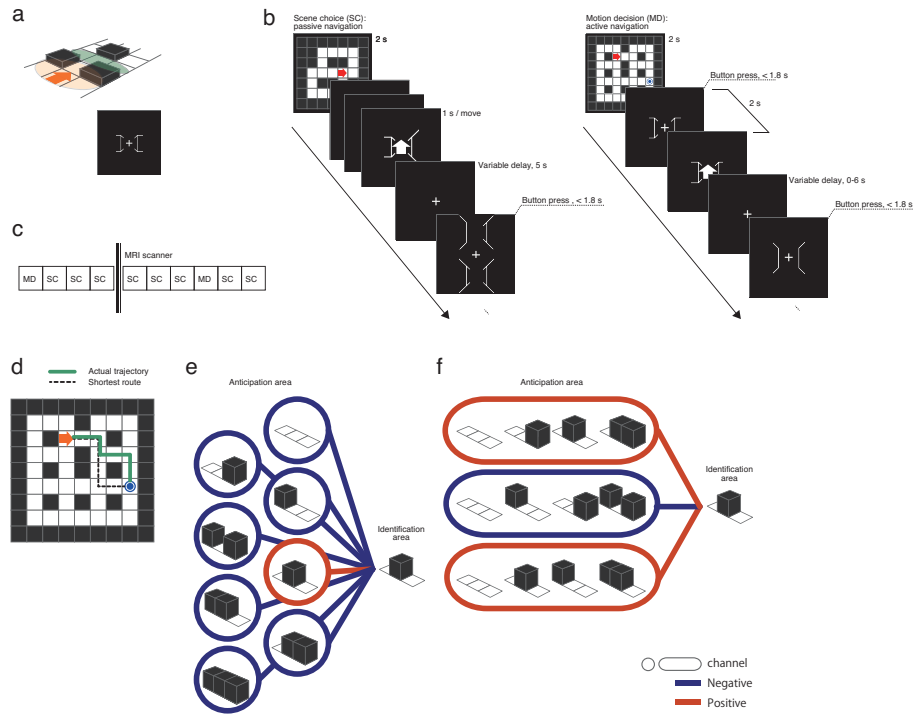


FIGURE 5.1: Experimental procedure and encoding scheme. (a) A schematic drawing of a single trial in my spatial navigation task. An orange oval indicates the visual field seen by a participant at the state indicated by the red arrow (state = location + orientation); seen are two walls (the forward-left and forward-right view parts) and one path (forward-center view part) (bottom). After moving forward, the participant sees three new view parts (green oval). (b) Participants took part in scene choice (SC) and motion decision (MD) navigation tasks. In SC (left panel), participants predicted the next scene view consisting of three unseen view parts (forward-left, forward-center, and forward-right) consisting of either wall (black square on the top display) or path (white square) elements. For each trial, they chose the next scene view from between the correct next scene and an incorrect one. In MD (right panel), participants were requested to navigate from an initial state (red arrow on the top display) to a goal square (blue circle). (c) Before the experiment, participants were trained in relating the 2D and 3D views in MD. (d) Sample behaviors in MD. Shown are the actual trajectory taken by the participant (green line) and the shortest route (black dashed line). This participant took the second-shortest route to the goal square. Route length was calculated as number of moves including pure rotations. (e) I assumed a perceptron architecture, in which multiple, different encoding channels cooperatively represent the next scene view. There were 8 possible scene views, so a naïve encoder design would have eight corresponding channels; when predicting a specific scene view, one channel is activated (red circle), while the others are inactivated (blue circles). (f) A more sophisticated encoder. The top and bottom channels (activated) each vote as expecting one of the four possibilities in each red oval, and the middle channel (inactivated) votes as expecting the remaining possibilities in the blue oval. No channel can predict the scene view by itself, but the majority of votes can (here, a predicted wall in the forward-center view).

3D scene views, each accompanied by a pre-indication of the next move direction. Visual feedback indicating the next direction was presented as a white arrow on the center of the screen. 'Up arrow' indicated forward movement, and 'left arrow' or 'right arrow' indicated left or right turns, respectively, while staying at the same square.

The first and second movements could be any of the three movement types (move forward or turn left/right), while the third movement was constrained to be forward. Before the third movement but after a pre-indication of the third move was presented, a delay period was introduced.

In the following choice period, participants were requested to choose as soon as possible, using an MRI-compatible button box, between two options, a correct scene, and a distractor scene, both represented in 3D wire-frame format. By their choice, participants were to indicate which scene they expected to see after the third movement. The new 3D scene after the third and forward movement contained three newly observed view parts (forward-left, forward-center, and forward-right), which had never been seen in that trial in the preceding move period or in the delay period.

Motion decision (MD) task

In the MD task, the participants navigated the same 9×9 map as was used in the previous practice task (Fig. 5.1d). A single block started with presentation on the map of the initial state and the goal position and ended with goal achievement. At the onset of a block, the initial state (position and body orientation) and the goal position were presented on a screen for 4 s. At each trial, a 3D wire-frame scene view at the current state was displayed and the participant was requested to make a decision as to how best to reach the goal, by pressing one of three action buttons, 'go forward', 'turn left', or 'turn right', of a three-key MRI-compatible button box, within a fixed decision period of 1.8 s, but as soon as possible. After the participant made a decision, visual feedback indicating the selected direction was presented as a white arrow on the center of the screen (feedback period, > 0.2 s). After a fixed time consisting of the decision period and the feedback period (2 s in total), a variable delay time of 0-6 s was imposed, then a new trial was started by displaying the next scene view.

After goal achievement, the next block was immediately started. While the combinations of initial states and goal positions were different between blocks and their orders were further different between participants, they were extracted from a predefined pool that was common over the participants. The MD session comprised a fixed number of 300 MR scans and admitted termination even in the middle of a block; only the data for completed blocks were used for analyses.

Since one of the eight participants did not complete both the SC and MD tasks, further analyses were performed only on the remaining seven participants.

5.2 Analysis methods

5.2.1 Behavioral analysis

All behavioral data analyses were performed with MATLAB (MathWorks) with the help of the Statistics and Machine Learning toolbox. When searching for the shortest route in each MD block, I used Dijkstra's algorithm implemented as a custom program in MATLAB. Reaction time (RT) analysis was performed after conversion of the RT data to within-subject Z-scores.

5.2.2 fMRI data preprocessing

The first six scans of each run were discarded so as not to be affected by initial field inhomogeneity. The acquired fMRI data underwent 3D motion correction using SPM12 (<http://www.fil.ion.ucl.ac.uk/spm>). The data were coregistered to the whole-head high-resolution anatomical image, and then spatially normalized to the MNI space (Montreal Neurological Institute, Canada). I then identified the 37,262 voxels of the cerebral cortex. The fMRI signals then underwent quadratic polynomial trend removal and noise reduction by means of singular-value decomposition ($K = 3$), and were temporally normalized within each session.

5.2.3 Encoding models

Whole-scene encoding model.

In my experimental setting, each scene view to be predicted could be represented by the path/wall status of three view parts, forward-left, forward-center, and forward-right, and thus represents one of eight possibilities ($2^3 = 8$). In any binary encoding scheme, each scene view out of the eight possibilities is represented as a code word, a string of binary codes ('0' or '1') of which each corresponds to a single channel (Fig. 5.1e). The simplest way to distinguish eight classes is to use an eight-bit representation, each bit designating one scene view (sometimes called one-versus-the-rest encoding; 1R); the code matrix Z^{1R} becomes the 8×8 identity matrix. The 1R encoding model X^{1R} ($[N \times 8]$, N = number of samples) is defined as $x_n^{1R} = z_{c_n}^{1R}$, the left-hand-side is the n -th row vector of X^{1R} and the right-hand-side is the c_n -th row vector (code word) of the code matrix Z^{1R} , where c_n is the index of the known scene view to be predicted.

View-part encoding model.

The view-part encoding model is the same as the original three-bit representation of the scene view. Then, the code matrix of the bit-wise encoding Z^{bit} is an 8×3 matrix. Every column has four '1's and four '0's (Fig. 5.1f). The view-part encoding model X^{bit} is an $N \times 3$ matrix, whose row vector is the bit-encoding code word of the scene view to be predicted for the corresponding data point.

Full encoding model.

Since each encoding channel assigns '1' to one or more scene views out of eight candidates and '0' to the others, there are in total 254 ($2^8 - 2$) possible meaningful channels, excluding all '0's or all '1's (i.e., all or none). For each scene view, a 254-dimensional vector consisting of the binary assignment over these 254 channels becomes its code word in the full encoding model. The code matrix of full encoding, Z^{full} , is then an 8×254 matrix. The full encoding model X^{full} is then represented as an $N \times 254$ matrix, whose row vector is the full-encoding code word of the scene view to be predicted for the corresponding data point. Full encoding includes all the encoding channels in the whole-scene model.

5.2.4 General linear model analysis

To perform data-driven analyses, which will be explained below, I estimated general linear models (GLMs) for the full encoding model, using as input the delay-period fMRI activities in sessions 1, 3, and 5 of the SC task. The GLM optimized the beta weights for those 254 channels to make them fit to the fMRI data over all the training trials, according to the multi-variate regression analysis. The code in each code word is either 0 or 1; if the code is 1, that channel is included in the regression of the fMRI activities, but if it is 0, the channel is not included. Although the number of explanatory beta values is fairly large, $254 \times V$ (V = number of fMRI voxels), it could be optimized because of the even larger number of fMRI activities. The regression parameters were determined by the L2-regularized least squares method, in which the regularization hyperparameter was tuned by applying a 10-fold cross-validation procedure to the training dataset, searching from 2 to 2^{20} with a logarithmically scaled interval. Six nuisance parameters that remove motion artifacts due to realignment were added to the GLMs. To avoid confounding any activation related to the visual presentation with button-press-related activity, no hemodynamic response function was convolved onto the fMRI signals.

When examining brain regions involved in prediction of upcoming scene views, I used another GLM whose design matrix corresponded to the whole-scene model (Fig. 5.7a; see below).

5.2.5 Data-driven analysis

Data-driven analyses on the different encoding models were performed by using the estimated GLM weight values. To examine the plausibility of different encoding models, I introduced two measures: the predictability of the fMRI BOLD signals based on the GLM weights, and the amount of error-correction ability attributable to the code matrix. GLM weights were further used to investigate the cortical correlates of scene prediction.

Predictability of encoding models

I examined the voxel-wise Pearson's correlation coefficient between measured fMRI signals and fMRI signals reconstructed by the GLM whose design matrix was defined as corresponding to the full encoding model (see Subsection 5.3.2). When displaying

the prediction accuracy (Fig. 5.7a), I used the 10th and 90th percentiles and median, rather than the mean, of voxel-wise correlation coefficients over the cortical voxels in each ROI, to reflect the intracortical localization of cerebral functions.

Figures 5.7a and 5.3 show the prediction accuracy for each fold in the 10-fold cross-validation (CV) when the GLM used a code matrix with different numbers of channels taken from the channel list of the full encoding model.

Appearance frequency matrix of scene view

I examined how robustly the set of code words in each encoding model were represented by reflecting the maze topography (see Subsection 5.3.2). When navigating a maze, information regarding scene anticipation should be transmitted more robustly for views that are more frequently seen in the maze versus those less frequently seen. I first constructed an appearance frequency matrix of scene view, i.e., an 8×8 matrix $\mathbf{a}$, whose element $a_{(c,c')}$ was the minimum of appearance frequency values of scene views c and c' . Here, the appearance frequency value was a normalized element value of the view histogram in Fig. 5.7b, left. This appearance frequency matrix does not depend on any encoding model. To measure the robustness of an encoding model, I calculated the Pearson's correlation between the off-diagonal part of the encoder's distance matrix (Fig. 5.7b, upper right) and the corresponding part of the appearance frequency matrix. A higher correlation implies the encoder's higher robustness against bit-inversion errors which reflects the map topography; if a particular scene is frequently seen in the maze, the corresponding code word should be isolated (then, error-tolerant), associated with larger element values in the distance matrix, while another scene that is less frequently seen can be associated with smaller element values in the distance matrix (less error-tolerant, but errors barely occur).

When examining the similarity between the appearance frequency matrix and the distance matrices of random encoding models, I prepared 1,000 random encoding models, each with 83 channels randomly taken from the full encoding channel list; the number of channels therein was fixed at the median number of the optimal encoding models over participants.

Robustness of encoding models

Consider an encoding process in which each encoding channel tries to transmit either of two symbols, '0' or '1'. Due to the presence of noise, or more precisely, bit-inversion error, a transmitted '0' may sometimes be received as '1', or vice versa. If there is no prior information available about these transmission channels, the most reasonable way to decode from multiple noisy channels is Hamming decoding, which selects the class whose correct code word is of minimum Hamming distance (Hamming distance = the number of bits different between the two binary code words) from the 'predicted' code word. In this study, the latter is defined as the vector aligning the outputs from binary classifiers, each corresponding to a single encoding channel (see Subsection 5.3.2). The results of this analysis are presented in Figures 5.2b, 5.4b and 5.5b.

Another criterion for the robustness of each encoding model is the capability of error correction. To quantify this, I applied every possible bit-inversion error to each of the eight code words and evaluated how well the Hamming decoding worked. When the Hamming decoding could lead to the right class, because the code word suffering from the bit inversion still fell at or below the minimum Hamming distance from the correct code word of the right class, the trial outcome was classified as 'accurate identification'. If decoding led to a wrong class, the outcome was classified as 'misidentification'. The remaining possibility, where the bit inversion made the noisy code word equidistant from correct code words of two or more classes, was called 'unidentifiable'. I can define these three categories even when each code word is disturbed by simultaneous inversions of two or three bits. If the proportion of the first category 'accurate identification' is large, the corresponding encoding model is said to have the characteristic of error correction. The results of this analysis are presented in Figs. 5.2c, 5.4c and 5.5c.

Characterizing information gains

To examine cortical localization of scene view information, I introduced an information gain index (IGI) based on an information gain (Quinlan, 1986). The information gain is the change in the information entropy, an average achieved by getting information through a single encoding channel (see Subsection 5.3.3). The information

gain $gain(q)$ of channel q was defined as

$$\begin{aligned} gain(q) &= 1 - I_q \\ I_q &= \frac{m_q}{8} \log_2 \frac{8}{m_q} + \frac{8 - m_q}{8} \log_2 \frac{8}{8 - m_q} \end{aligned}$$

I_q quantifies the information (in bits) that would be obtained by observing a single channel q , if the observation is noiseless, and m_q is the number of code '1's in this channel q . If there is no information regarding a single view part, the information entropy is 1 ($= \log_2 2$), because there are two equal possibilities (path/wall). The information gain represents how much this initial entropy is reduced, so it falls within $[0, 1]$: 1 indicates that the channel provides certain information (without ambiguity) of the status of the view part; 0 indicates that the channel provides no information.

For the three view parts, forward-left (L), forward-center (C), and forward-right (R), I calculated the information gain $gain(q)$, represented as $gain(q)_L$, $gain(q)_C$, and $gain(q)_R$, respectively. To comprehensively represent the information gain over all possible channels, I computed an information gain index (IGI) that quantifies the overall contribution of a specific cortical voxel to the prediction of each view part:

$$IGI_i = \left[\sum_{q=1}^{254} w_{q,i}^{full} gain(q)_L, \sum_{q=1}^{254} w_{q,i}^{full} gain(q)_C, \sum_{q=1}^{254} w_{q,i}^{full} gain(q)_R \right]$$

where w^{full} is the absolute weight of the GLM constructed from the full encoding model and i is the voxel index.

When mapping onto the cortical surface, I applied voxel-wise t-tests. The null hypothesis was that the IGI comes from a normal distribution with a mean of 0.13 (mean of the above-defined gain over channels) times a baseline activity (mean activity of all scans), and with unknown variance. Figure 5.7c shows the significant voxels in terms of the IGI ($p < 0.1$).

5.2.6 Decoding analysis method

Training of scene anticipation classifiers

For each channel of an encoding model, a binary classifier called a linear support vector machine (SVM) was trained, using the SC sessions 1, 2, 3, and 5 over the whole cortex. Based on the fMRI voxel-wise signals and the set of code words defined by the encoding model, a set of SVM classifiers (whose number is the same as the number of encoding channels) was trained, to be ready for decoding. Given input voxel values $y = [y_1, \dots, y_v]$ (V is the number of voxels), SVM provided a discriminant value for classifying between scene view classes assigned a label '1' and those assigned a label '0', based on the code matrix of the encoding model. Considering the relatively small number of samples available for training and testing, I did not tune the 'C' constant in the SVM ($C = Inf.$), as previously described (Reddy, Tsuchiya, and Serre, 2010). It should be noted that each SVM classifier was invariant in its function over exchange of the labels '0' and '1', and that the Hamming decoding was invariant over permutation of channel order. Thus, in the decoding analysis, I redefined the full encoding model to remove unnecessary channels; when selecting one channel from the set of equivalent ones, I selected the earlier one on the channel list aligned with respect to GLM weights (see subsection 5.3.2). Additionally, I introduced a subsampling scheme for removing imbalances in the number of training samples assigned '1' (positive) and '0' (negative); simple use of an unbalanced dataset to train binary classifiers may introduce a bias into the classifiers. Let N_{pos} and N_{neg} be the numbers of positive and negative samples in the original training dataset, respectively. If $N_{neg} > N_{pos}$, then N_{pos} negative samples were randomly subsampled, otherwise N_{neg} positive samples were subsampled, so that the numbers of positive and negative samples in the training dataset became equal.

When decoding the next scene view from the set of SVM binary classifiers, I used a Hamming decoder. According to Hamming decoding, the output is the one class out of the eight possibilities whose code word (which depends on the encoding method) has the smallest Hamming distance from the predicted code word, the latter being the set of binary outputs from the SVM classifiers. It should be noted I did not use any prior information such as map topography to supplement the participant's brain activities.

Scene reconstruction

When visualizing the predicted scene view during active navigation (see Subsection 5.3.4), I first calculated the wall probability vector in the scan i , $P(v_i)$, based on a vote of the outputs from the SVM classifiers.

$$\begin{aligned} P(v_i) &= \sum_{c=1}^8 P(v_i = c) \times v^c \\ v^c &= [\delta^L(x_c), \delta^C(x_c), \delta^R(x_c)] \\ \delta(x) &= \begin{cases} 0, & x = \text{path} \\ 1, & x = \text{wall} \end{cases} \\ p(v_i = c) &= \frac{\sum_{q=1}^Q \hat{x}_{i,q} \equiv z_c}{Q} \end{aligned}$$

where Q is the number of encoding channels (same as the code length). Then each element $P(v_i)$ had an analogue value, ranging between 0 and 1. To regulate the inhomogeneity of basal brain activities, I binarized each element of $P(v_i)$ to obtain the predicted scene view (path/wall) by applying a block-dependent threshold set to the within-block median of $P(v_i)$. Figure 5.11 shows the results.

Time-shift Hamming distance

To investigate the decodability of two or more steps of anticipation (see Subsection 5.3.4), I defined the time-shift Hamming distance (TSHD) D as the Hamming distance between the predicted code word $\hat{x}$, which is the alignment of the SVM's outputs based on current brain activity, and the code word x for a correct, possibly future, scene view with a lag time j :

$$D_{x\hat{x}} \begin{cases} \frac{\sum N - jn = 1 d_h(x_{n+j}, \hat{x}_n)}{N - j}, & j \geq 0 \\ \frac{\sum Nn = 1 - j d_h(x_{n+j}, \hat{x}_n)}{N + j}, & j < 0 \end{cases}$$

where $d_h(x, y)$ is the Hamming distance between x and y . Although I was mostly interested in decoding future scene views, I also examined the decodability of past

scene views; in the past case, the lag time j was just a negative integer. The results are shown in Figure 5.11 and Fig. 5.12.

5.3 Results

5.3.1 Behavioral results for passive and active navigation

Seven participants completed both the SC task and the MD task in an fMRI scanner, during which the participants saw 3D wire-frame views presenting egocentric scenes constructed of open paths and walls (Figs. 5.1a–c).

In SC, trained participants were requested to choose from two options which scene would appear after a move selected by the computer: the correct upcoming scene or a distractor scene (Fig. 5.1b, left). I found that participants were able to choose the future scene accurately. Quantitatively, the results were: $93.4\% \pm 2.6\%$ correct, $5.6\% \pm 2.3\%$ incorrect, and $1.0\% \pm 1.0\%$ missed (mean $\pm$ SEM across participants; chance accuracy = 50%). See Chapter 4 for further details. Incorrect trials were excluded from subsequent analysis, as were missed trials, in which participants did not press the answer button in the allotted time (1.8 s).

In MD, the participants steered to an instructed goal from an initial state over a series of trials, each trial requiring a choice of one of three decision options: move forward, turn left, or turn right (Fig. 5.1b, right). Experimental blocks were separated by arrival at instructed goal positions. The number of completed blocks was 11.43 ± 0.98 (mean $\pm$ SEM across participants), while the uncompleted blocks (final block in each session) were excluded from subsequent analysis. I found that participants could accurately trace the shortest or the second-shortest route (i.e. one or two moves longer than the shortest route, Fig. 5.1d) in the majority of blocks; the shortest route was traced in $83.7\% \pm 24.5\%$ of blocks (mean $\pm$ SD across participants), and the second-shortest route in $9.6\% \pm 18.3\%$ of blocks. Reflecting the topography of the MD map, many of the participants' decisions were forward moves ($75.9\% \pm 9.5\%$); left and right turns accounted for $12.7\% \pm 10.0\%$ and $11.5\% \pm 9.0\%$ of the decisions, respectively. RT did not differ among these decision types (Friedman test, $p = 0.07$).

5.3.2 Comparison between encoding models

Next, I sought the encoding model that best encodes the participants' prediction of the upcoming scene view into their voxel-wise fMRI activity, which was represented by the delay-period activity in the SC task. There were eight prediction possibilities, since each of the three view parts, forward-left, forward-center, and forward-right, could show only two options (path or wall). Each encoding channel might thus assign '1' (positive) to one or more out of the eight possibilities, and '0' (negative) to the others (Figs. 5.1e, f); hence, there could be 254 ($= 2^8 - 2$, all and none were excluded) channels in total, which constitutes the full encoding model (see Subsection 5.2.3). When these channels were aligned (sorted in descending order) according to their GLM weight value, I found that the multivariate prediction accuracy saturated after incrementally incorporating about 75 aligned channels (Fig. 5.2a and Fig. 5.3, see Subsection 5.2.4). This implies that the fMRI signals were expressible by using a smaller number of features, which would be less costly than the full (most expensive) representation.

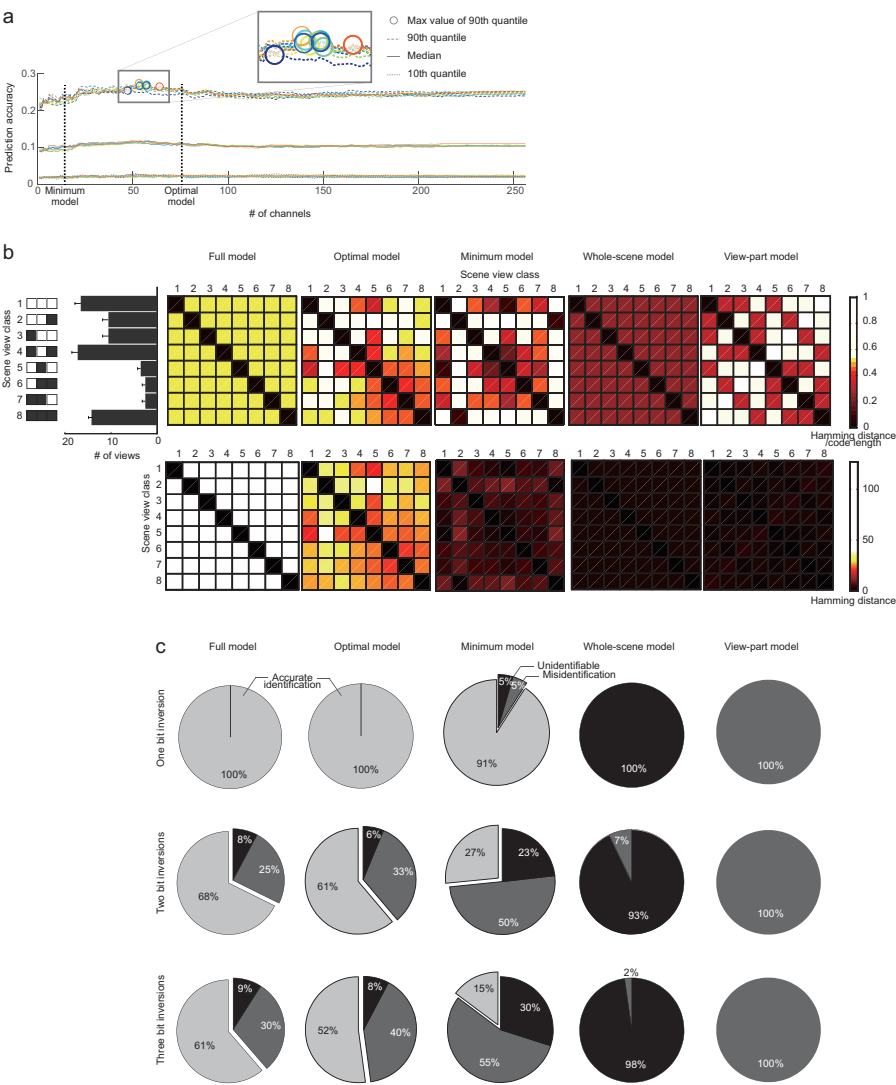


FIGURE 5.2: Characteristics of encoding models. (a) Voxel-wise prediction accuracy (in terms of Pearson’s correlation coefficient between measured and predicted brain activity. See Subsection 5.2.5.) became saturated with 40–70 encoding channels, and did not increase further with larger numbers of channels. A single line corresponds to the prediction accuracy for one validation set from ten folds in the 10-fold CV procedure from a typical individual (participant 5, bilateral IPG). The results of other individuals and ROIs are presented in Fig. 5.3. (b) Distance matrix in terms of the Hamming distance between pairs of code words of the eight view classes (depicted in the left-most inset; white, path; black, wall), under full encoding (‘Full model’), data-driven optimal encoding (‘Optimal model’), data-driven minimum encoding (‘Minimum model’), naïve eight-class encoding (‘Whole-scene model’), and view-part-wise encoding (‘View-part model’). Top panels indicate the Hamming distance normalized by code length and bottom panels are for the original Hamming distance. The histogram beside the full encoding model shows the normalized frequency of the eight scene views in the SC task. Error bars indicate SD over three different maps. My data-driven models reflected the scene view frequency, a characteristic of map topography, such that more frequently seen views were more frequently predicted. (c) Full encoding and optimal encoding are robust against errors introduced when observing the channels, such that a one bit-inversion error on every encoding channel can be corrected (first row, first and second columns). On the other hand, any single bit-inversion causes unidentifiable codes in whole-scene encoding (fourth column), and leads to misidentification in view-part encoding (fifth column). The second and third rows show the situation when two and three bit-inversion errors occur, respectively. For the detailed method used to calculate these pie charts, see Subsection 5.2.5, (b, c). For the other participants’ data, see Figs. 5.4 and 5.5.

My finding that the fMRI signals in SC were well predicted using a small number of encoding channels led to the idea of a *minimum encoding model*. By incrementally incorporating single channels from the aligned channel list sufficient for distinguishing the eight possibilities, I defined the minimum encoding model (see Figs. 5.4a and 5.5a). As described in the above paragraph, each channel separates the eight scene views into positive and negative groups. As the number of channels increases, the number of identifiable scene views would also increase, as expected. I included channels one at a time, according to the aligned channel list, and stopped adding channels when the number of identifiable views reached eight; this procedure created the minimum encoding model. When using a randomly permuted channel list, the number of channels constituting the minimum encoding model was 6.14 ± 1.68 (the SD was calculated over 5,000 random permutations). However, using my incremental addition paradigm, the smallest number (11.43 ± 2.76 ; mean $\pm$ SD over

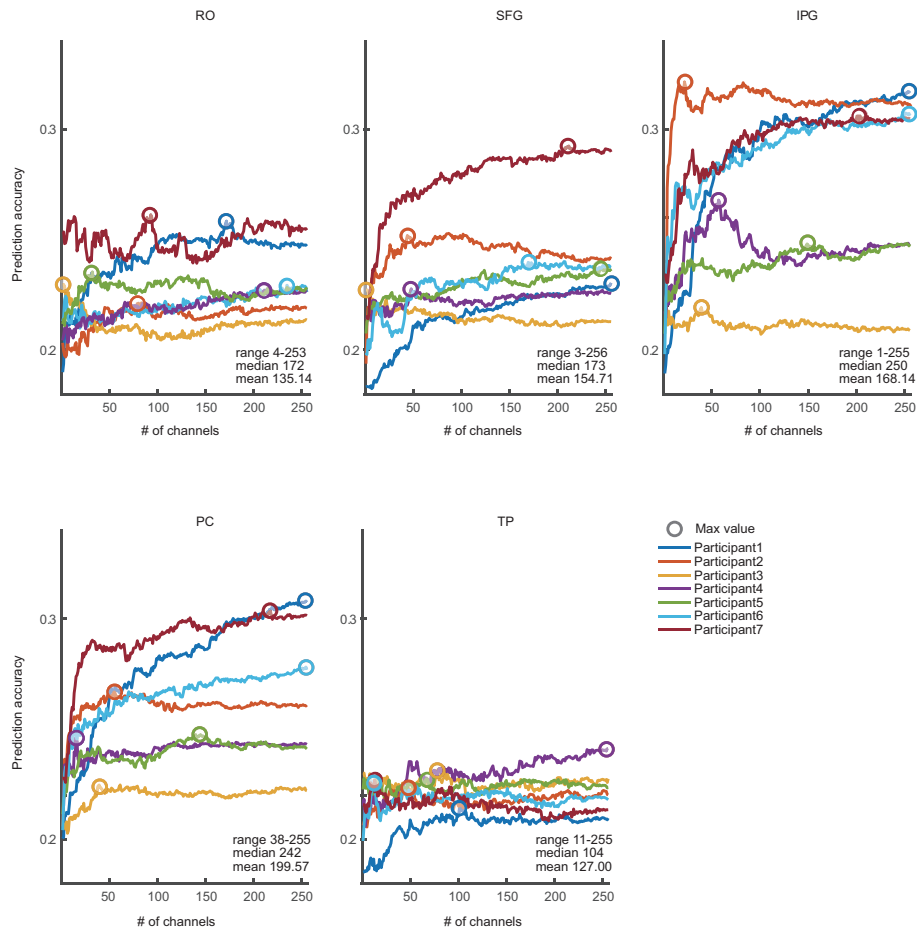


FIGURE 5.3: Averaged prediction accuracy over ten folds in the 10-fold CV in each individual. Each line indicates the mean prediction accuracy of the top 10th percentile in each region of interest (ROI). The range, median and mean (over individuals) channel numbers needed for maximum prediction accuracy when using the top 10% of ranked voxels in each ROI are shown as the right-bottom legend.

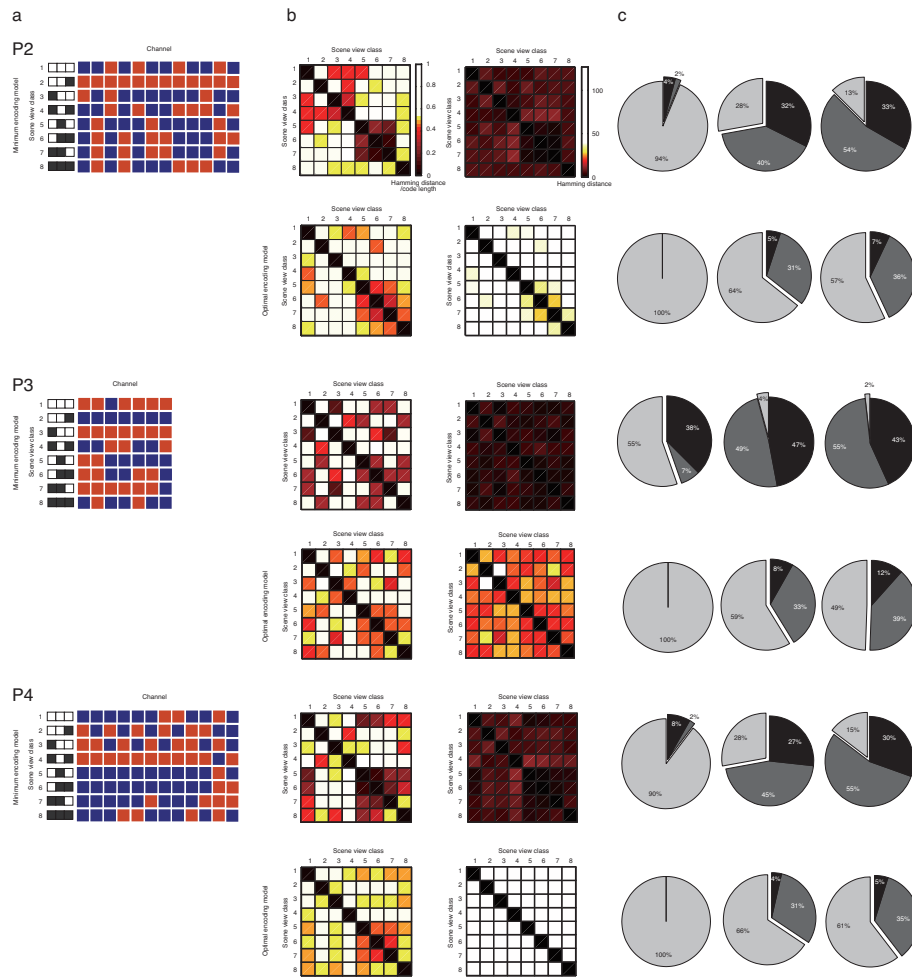


FIGURE 5.4: (a) Minimum encoding model, participants 2-6. The data-driven analysis obtained the minimum and optimal encoding models for each of participants 2-4, respectively. Left black-white matrix represents the eight types of scene view (white, path; black, wall). Right red-blue matrix represents the code matrix (red, activation, [1]; blue, control, [0]). Each row, i.e., a code word, in the code matrix corresponds to the scene view on the left. (b) A distance matrix consisting of the Hamming distances between eight scene view classes (top: minimum encoding model, bottom: optimal encoding model). In the left panel, the distance is scaled by the code length, but not in the right panel. (c) A pie chart showing how bit-inversion error will affect Hamming decoding. Light gray, dark gray, and black indicate, respectively, the frequencies with which bit-inversion errors will lead to correct decoding ('accurate identification'), incorrect decoding ('misidentification'), and ambiguous cases in which more than two classes become decoding candidates ('unidentifiable'). Left panels show the case of one bit-inversion error, middle panels the case of two simultaneous bit-inversion errors, and right panels the case of three simultaneous bit-inversion errors.

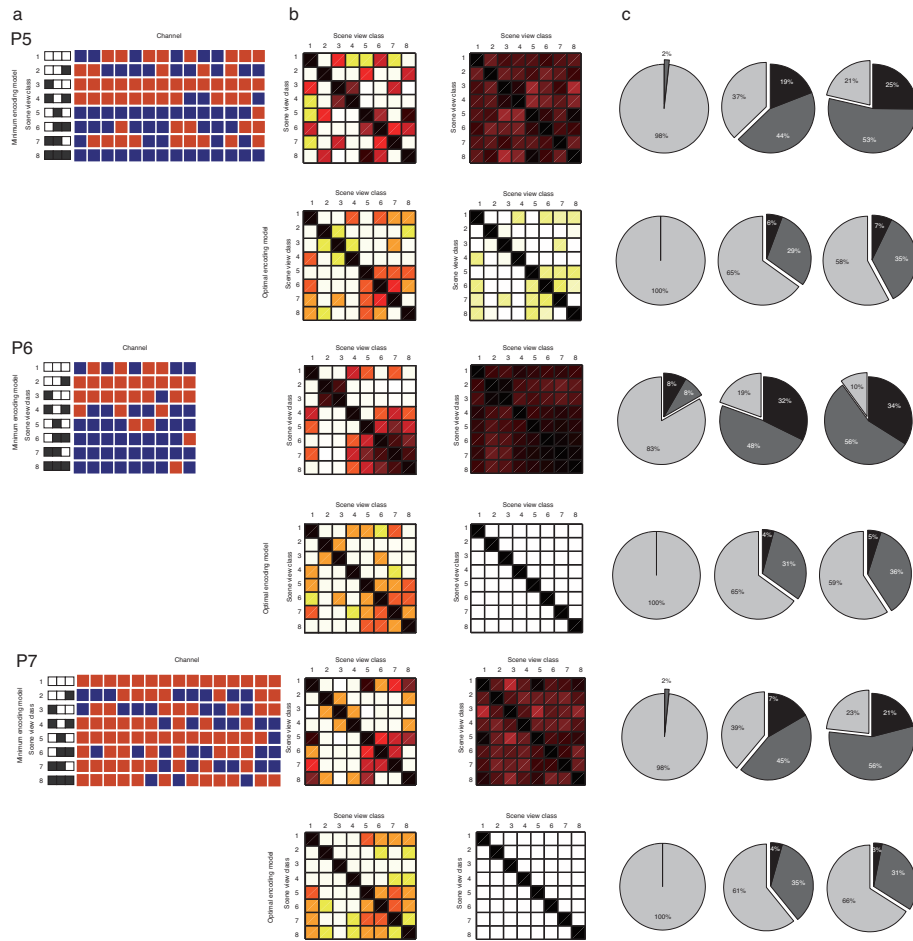


FIGURE 5.5: (a) Minimum encoding model, participant 5-7.

participants; range 7-15) was larger than when the channels were selected from the randomly permuted channel list.

Any encoding model would assign a binary code word to each of the eight possible views to be predicted; this code is expected to constitute the outputs of the encoding channels. Examining the Hamming distance (number of different bits) between each code word pair of the minimum encoding model in the form of a distance matrix (Fig. 5.2b, middle), I found that many code words of my minimum encoding model were represented in an idiosyncratic manner. For example, the third class [*wall, path, path*] is distant from its nearest neighbor (the fifth class [*path, wall, path*]), with a Hamming distance of four. Such an isolated class can enjoy error correction in decoding; indeed, even if one or two channel(s) are disturbed by bit-inversion error, the class isolation allows the Hamming decoder to decode this view class accurately. On the other hand, the fifth class [*path, wall, path*] has the first class [*path, path, path*] as its nearest neighbor with a Hamming distance of one. Although one bit-inversion could thus lead to misclassification in this case, such misclassification occurred only infrequently during navigation, because of the infrequent occurrence of the wall status in the forward view [**, wall, **] in the environment, which in this case was the maze topography in SC (Fig. 5.2b), my minimum encoder realized robustness in its decoding in a data-driven manner. On the contrary, the whole-scene encoding (grandmother cell-like) model, which is the simplest decomposition of a multi-class classification problem into multiple binary classification problems, includes error detection among its characteristics, but cannot correct bit-inversion errors (Figs. 5.2c, 5.4 and 5.5). It spaces every code word in an equidistant manner, without reflecting any information from the task environment. The most naïve encoding model, the view-part model, provides the most efficient representation of the scene views, but has neither error detection nor error correction ability (Fig. 5.6); a single bit-inversion error always leads to misclassification.

Although the minimal encoding model had the minimal number of channels, its error correction ability was not sufficient. When I increased the number of channels, I found not only that the prediction ability of the fMRI signals became stable (Fig. 5.2a), but also that the decoding ability of the Hamming decoder increased (Fig. 5.9a). I called this expanded encoding model the optimal encoding model (see Subsection 5.2.6). My data-driven encoding models reflected the map topography-dependent frequency of the eight scene views (histogram in Fig. 5.2b, see Subsection 5.2.5)

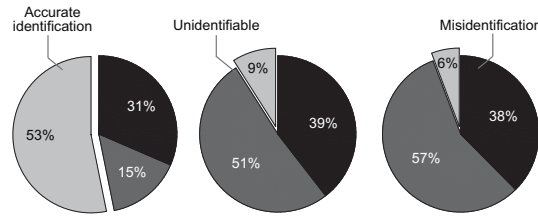


FIGURE 5.6: Effects of bit-inversion errors in the data-unrelated minimum encoding model. Each value was calculated over 1,000 random permutations of the full channel list.

better than data-unrelated, randomly assigned encoding models (optimal encoding model, Pearson's correlation coefficient = median across participants 0.21, range 0.05-0.56; data-unrelated random encoding model, median -0.01, range -0.61-0.61. The Wilcoxon rank sum one-sided test, $p < 0.05 \times 10^{-1}$).

5.3.3 Scene anticipation in frontal and parietal cortices

To identify the cortical regions involved in scene anticipation, I examined voxel-wise weight values in another GLM that reproduced the fMRI BOLD activities with the whole-scene model. I found substantially large voxel-wise absolute weights consistently across participants in five bilateral brain regions: rolandic operculum, superior prefrontal gyrus (SFG), inferior parietal gyrus (IPG) and precuneus (PC) and temporal pole (TP)(Fig. 5.7a). Although the whole-scene model was used here because the encoding channels in this model were mutually orthogonal, I verified that the above results remained unchanged with use of the full encoding model.

Previous neurophysiological studies have suggested that widespread brain regions, including higher-order fronto-parietal areas, are involved in prediction (Doll, Simon, and Daw, 2012). When I examined the prediction-relevant voxels, defined as those which exhibited high correlations between measured and predicted brain activity consistently over participants, they were mainly found in the IPG and the PC (Fig. 5.7b). Figure 5.7c and 5.8 show the information gain index (IGI, see Subsection 5.2.5) mapped onto the cortical surface, in which a voxel is colored when it shows a significant weight value selectively for one of three view parts (forward-left, forward-center, or forward-right). The SFG and the TP were found to include a large proportion of voxels that showed high information gain irrespective of view part. This result suggested a functional difference in scene prediction between the parietal

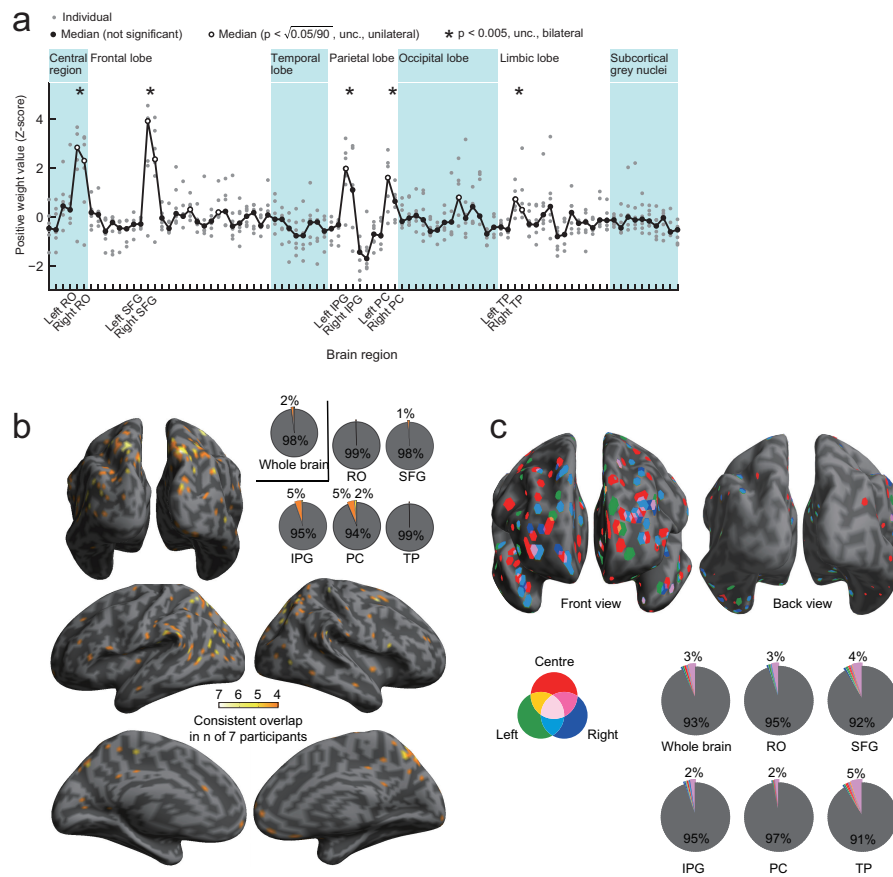


FIGURE 5.7: Brain regions involved in scene prediction. RO, rolandic operculum; SFG, superior frontal gyrus; IPG, inferior parietal gyrus; PC, precuneus. (a) Mean absolute weight value in the naïve eight-class encoding (whole-scene model, see Fig. 5.1e). The weight value was normalized across regions of interest (ROIs) within each participant into a Z-score with mean of 0.0 and standard deviation of 1.0. The abscissa indicates brain ROIs according to a list identified by automated anatomical labeling (Tzourio-Mazoyer et al., 2002). Each small, gray dot corresponds to a single participant, and each black, large one to the median of all participants. (b) Predictable-voxel maps showing overlap of the voxels consistently involved in the whole-scene model, plotted on the inflated brain surface. Bright parts consist of voxels involved in scene prediction in the full encoding model in at least 4 out of 7 participants; a statistical significance threshold of $p < 0.05$ (unc.) ($r > 0.21$), was required in each participant. The pie charts show the rates of bright-colored voxels in the respective brain ROIs. (c) Spatial distributions of the voxels contributing to decoding scene predictions. Colored voxels show those of substantial information gain index (IGI, for definition, see Subsection 5.2.5) at least 3 out of 7 participants and hence those that would be incorporated into the scene prediction process; the top 20% of voxels in terms of the IGI are plotted, whose colors correspond, respectively, to view parts (forward-left: ‘Left’, forward-center: ‘Center’ and forward-right: ‘Right’). The pie charts show the rates of colored voxels in respective brain ROIs.

regions and the SFG and TP; the IPG and PC were involved in scene anticipation itself, while the SFG and TP were specialized for the prediction of specific view parts.

5.3.4 Decoding results

For each channel constituting any encoding model, I trained a linear support vector machine (SVM) to output positive ('1') or negative ('0') based on the delay-period brain activities in SC. As the number of encoding channels increased, the accuracy of the Hamming decoder (see Subsection 5.2.6) was likely increased (Figs. 5.9a and 5.10), in accordance with the Condorcet jury theorem (Boland, 1989). However, the full encoding model employing all 127 channels ($= 254/2$, considering the symmetric nature of SVMs) did not show the highest decoding accuracy. The minimum encoding model did not show sufficient decoding accuracy either, due possibly to a lack of the necessary features (encoding channels). As a moderate choice, I defined an optimal encoding model that showed the highest decoding accuracy in SC, which was in turn evaluated using the MD task (cross-task validation). The decoding accuracy of this optimal encoding model was comparable to or even better than that of the full encoding model in MD (paired Wilcoxon signed ranks test, $p = 0.88$) (Fig. 5.9b). I also found that the decoding performances of the whole-scene model (eight channels) and the view-part model (three channels) were almost zero and at chance level, respectively. The decoding accuracy in MD of the optimal encoding model was significantly higher than chance and that of the view-part model (Wilcoxon signed ranks test, $p < 0.05$). Accordingly, the whole-scene model and the view-part model showed no cross-task decodability for scene view prediction, while my optimal encoding model allowed me to decode the upcoming scene views based on fMRI brain activities. The model was able to fully utilize the brain signatures of scene prediction during navigation, found to be represented by multiple channels in a distributed manner.

Next, I divided all the MD scans into those in which participants pressed a button (decision scans) and all others (delay scans). Examining the relationship between decoding outcomes based on fMRI activities during decision scans and the RTs in the trials covering the decision scans, I found that participants showed shorter RTs when the decoding outcome for the scene prediction was more successful than those of all MD trials (Fig. 5.9c). Moreover, this reduction in RT was most prominent

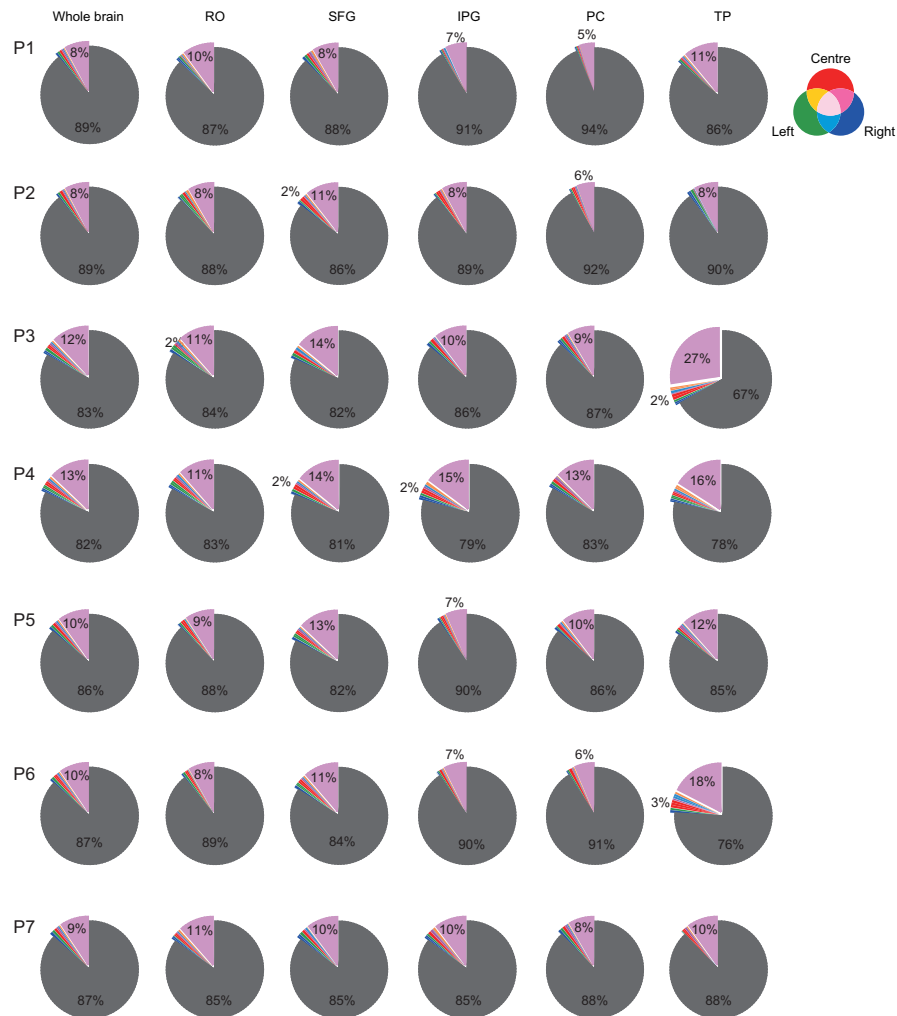


FIGURE 5.8: Individual results of spatial distributions of the IGI. Each row corresponds to a single participant (P1-P7).

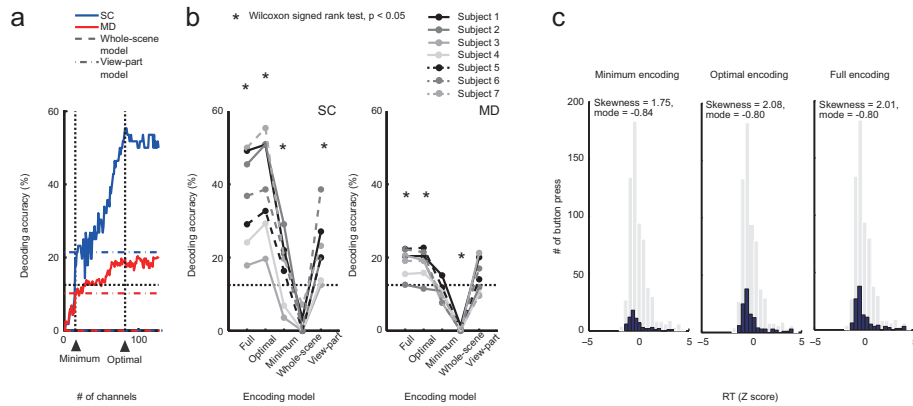


FIGURE 5.9: Decoding results. (a) Decoding accuracy for an individual (participant 1) against number of encoding channels progressively incorporated into the calculation, starting at the beginning of a list of channels sorted by the GLM weight value. The blue line shows the decoding accuracy validated in the SC sessions (sessions 1, 2, 3, 5, for training, and session 4 for test), and the red line shows the cross-task decoding accuracy for the MD task (SC sessions for training, and the MD session for test). Dashed lines and dot-and-dash lines indicate decoding accuracies with the whole-scene model and the view-part model, respectively. (b) Different encoding schemes lead to different decoding performances in the SC task (left), where each decoding accuracy was examined in the validation session (session 4). Similar tendencies were also observed in the MD task (right). Black asterisks indicate significance (Wilcoxon signed rank test, $p < 0.05$). (a, b) Horizontal, black dotted lines indicate chance level (eight classes, 12.5%). (c) When the Hamming decoder was successful in decoding the scene prediction, the participants showed shorter RTs (blue histogram) than those in all decoding trials in MD (gray, skewness = 0.90, mode = -0.63). This RT reduction was most prominent in the optimal encoding model. The RTs were normalized to have mean zero and variance one within each participant.

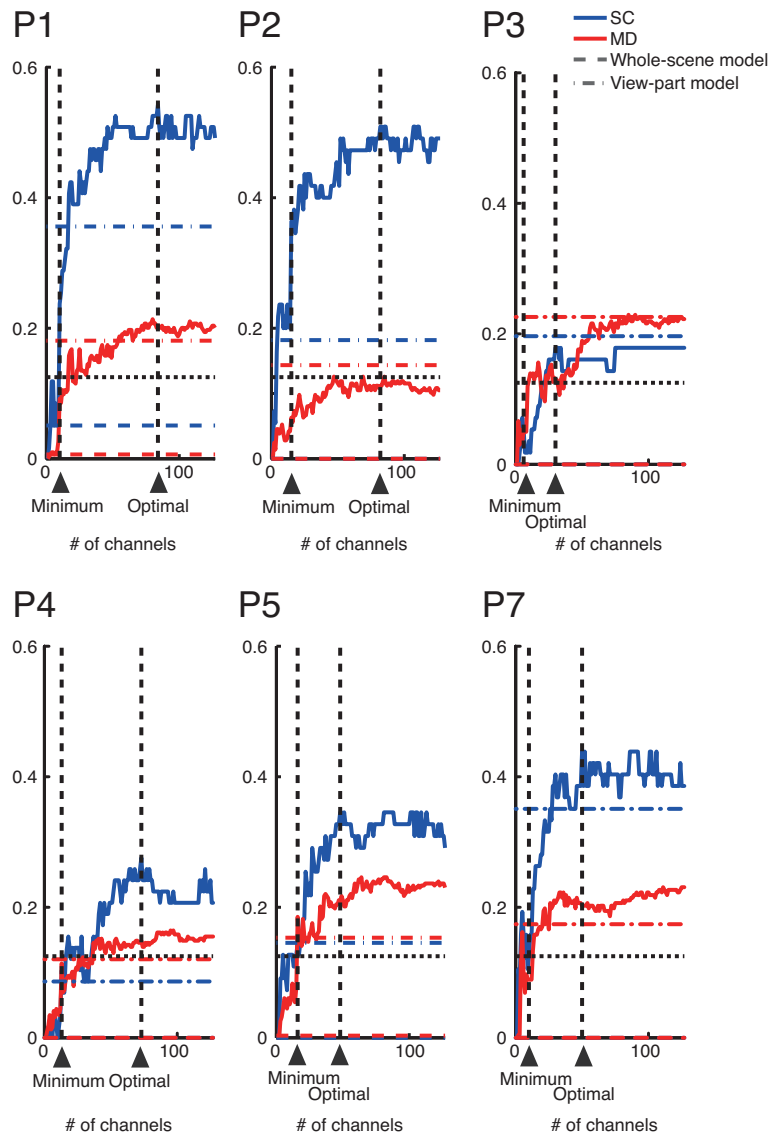


FIGURE 5.10: Decoding accuracy for participants 2-7, against number of encoding channels. Black arrows and vertical dashed lines, numbers of channels in the indicated models; horizontal, black dotted line, chance level. The blue line shows the decoding accuracy validated in the SC sessions (sessions 1, 2, 3, 5, for training, and session 4 for test), and the red line shows the cross-task decoding accuracy for the MD task (SC sessions for training, and the MD session for test). Colored, dashed lines and dot-and-dash lines indicate decoding accuracies with the whole-scene model and the view-part model, respectively. The former is often close to zero.

when employing the optimal encoding model. On the other hand, there was no significant difference in the decoding accuracy between the decision scans and the delay scans, regardless of the employed encoding models (Wilcoxon matched-pair signed ranks test: minimum model, $p = 0.81$; optimal model, $p = 0.58$; full model, $p = 0.94$). Accordingly, there was a tendency for the decoding to be successful when the decision RT was relatively short.

Although the participants had been requested to choose the correct next scene, making prediction of the next scene was crucial to the SC task, this was not necessarily the case in the MD task, because the MD participants were requested to reach their destinations with as small a number of trials as possible. Actually, the participants showed good planning performance in MD (see Subsection 5.3.1). Thus, anticipation was not limited to one-step prediction. To investigate the decodability of two or more steps of anticipation, I examined a time-shifted Hamming distance (TSHD), which measures the distance between the predicted code word based on current brain activity and the correct code word for a future or past scene view with some lag time. Interestingly, the TSHD did not show a minimum value with a lag time of 0 (one-step prediction), that is, smaller was not always better (see Figs. 5.11 and 5.12). In particular, when the participant's behavior was very successful, the TSHD showed minimum values for future scene views that were close to the goal (Figs. 5.11c). On the other hand, the blocks with a detour trajectory (defined as a trial where the number of steps exceeded the shortest route by 10 or more steps) had a larger TSHD between decoded code words in early steps, and those of real scene views around the goal (Fig. 5.11d).

5.4 Discussion

In my encoding/decoding analyses of fMRI activities during the performance of two types of 3D navigation games, I made two major observations. (i) Although my data-driven minimum and optimal encoding models were defined so as to exclusively and accurately reproduce the participants' brain activities, they also identified encoding channels that were effective in robustly transmitting scene anticipation. These encoding models allowed the simple Hamming decoder to work well even when the channels suffered from noise, because they reflected the characteristics of the navigation environment. They also allowed me to examine specific brain regions that

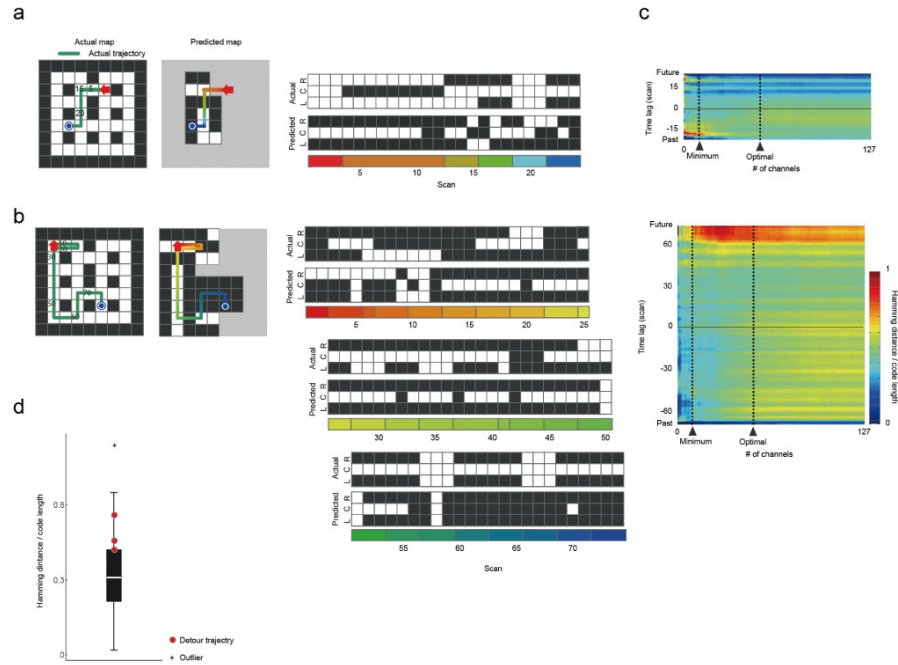


FIGURE 5.11: Sequential decoding in two MD blocks. (a, b) Scene reconstruction from single-scan fMRI activity (2 s) over a single MD block, but before this participant (participant 1) saw the actual scene. The left panel indicates the trajectory that was taken by the participant, starting from an initial state (red arrow) to a goal (blue circle). A digit on the map shows the scan index corresponding to the horizontal axis in the middle panel. The middle panels indicate the series of correct scene views (upper) and predicted scene views (lower). The hemodynamic response delay was not compensated for in this decoding analysis. L, forward-left; C, forward-center; R, forward-right. (c) TSHD, which signify the discrepancy between a real scene (not necessarily the present one) and the predicted one decoded from current brain activity. The top and bottom panels correspond to the MD blocks in (a) and (b), respectively. (d) Distance between the code words predicted for the scenes around the goal position based on the initial three scans and the code words of the corresponding true scenes. The optimal encoding model was used. In the blocks with detour trajectories, the distance was longer than the median of the other trajectories (white line). The bottom and top edges of the box indicate the 25th and 75th percentiles, respectively. A point located beyond 1.5 times the box height was defined as an outlier.

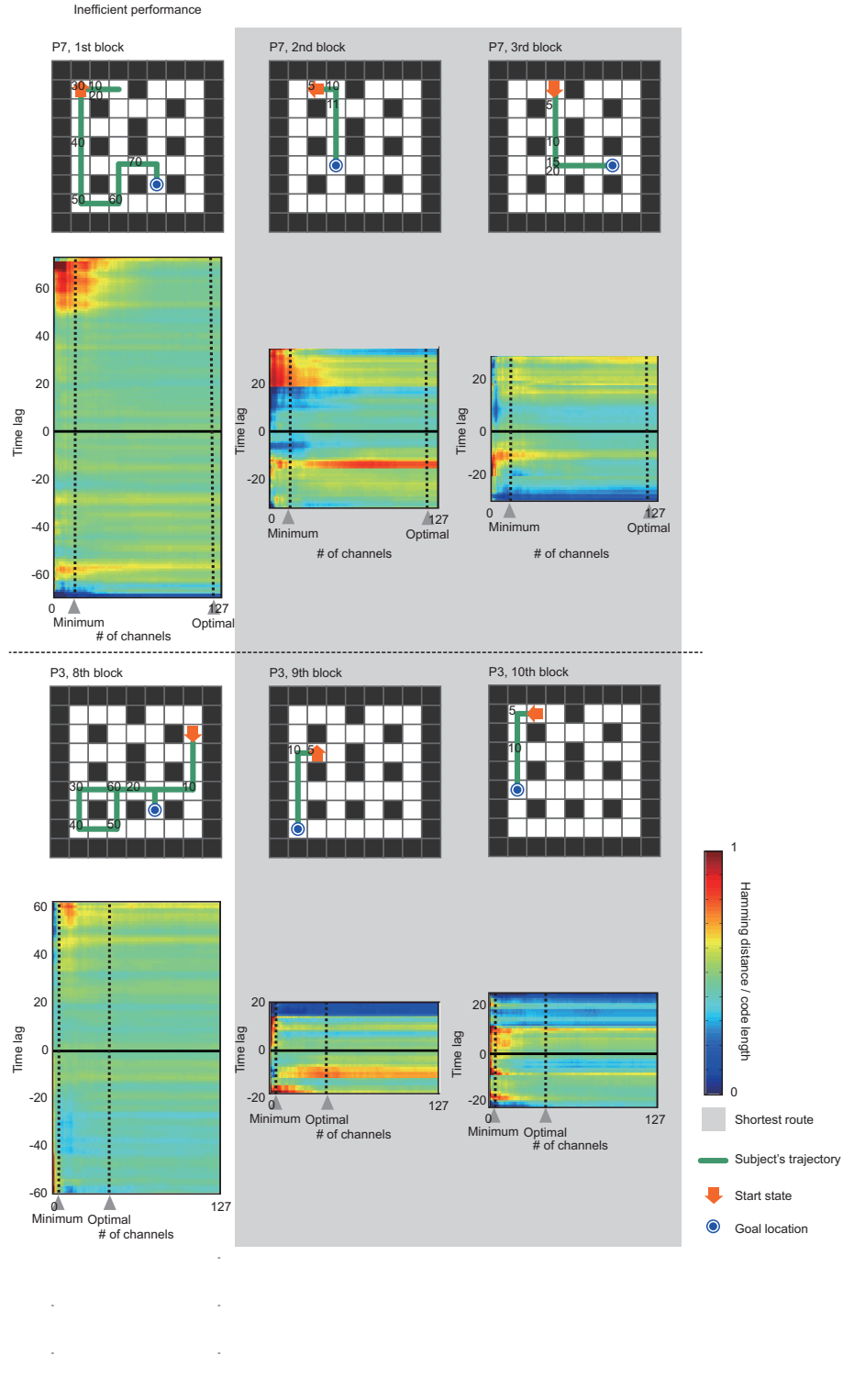


FIGURE 5.12: Action history and TSHD during the motion decision task. The TSHD between the code word decoded from current brain activity and one of the correct code words (of the future or past view, being dependent on the sign of the time lag, or the next view when the time lag is zero). Shown are results for three trials, one a case where the participant took a detour resulting in a trial score of 10 or more greater than that of the shortest route (left panel) and two cases where the participant took the shortest route (middle and right panels). Upper and bottom panels are for participant 7 and participant 3, respectively.

can be involved in those encoding models. (ii) By using such data-driven encoding models, I realized not only better decoding ability than the naïve encoding models, but also reasonably good cross-task decodability. Actually, my encoder/decoder pair was obtained only from data recorded during passive navigation (i.e., the SC task), but showed significant decoding ability even when applied to the data recorded during active navigation (i.e., the MD task).

The encoding strategy of the brain, especially for its high-dimensional cognitive states, is comparable to the problem of feature representation in the field of machine learning. Error-correcting output coding (ECOC) provides a general framework for representing high-dimensional features in a robust manner (Dietterich and Bakiri, 1995; Yukinawa et al., 2009; Takenouchi and Ishii, 2009). According to ECOC, an original multi-class classification problem is decomposed into multiple binary classification problems; each of these then acts as a transmission channel, transmitting the message of the original classes, even if it may be noisy. When using a larger number of transmission channels than the number of original multiple classes, a receiver can detect an original message based on received codes that have passed through noisy transmission channels. Thus, the ECOC provides a powerful encoding scheme to allow robust information transmission through multiple but unreliable channels. In this study, I presented data-driven encoding models based on ECOC. There are three major reasons for the plausibility of ECOC-based encoding models in the brain. First, neural information pathways (e.g., axons) are probabilistic rather than deterministic, due to the stochastic nature of neuronal spikes and the instability of axonal transmission (Bakkum et al., 2013). Thus, the neural decoding system inevitably needs robustness against the probabilistic factors (noise) involved in its transmission system. Second, since each neural element (neuron) can carry binary information (spike or non-spike), the network state at a single moment is of a binary vector. One natural scenario to solve a multi-class problem by using such a binary system is to decompose the multi-class problem into multiple binary problems, i.e., ECOC. Third, the brain has massively parallel processing machinery (Pashler, 1994; Sigman and Dehaene, 2008), leading to the idea that its encoder could consist of massively parallel, but noisy, channels. On the other hand, in ECOC, there is a trade-off between communication cost and ability to correct errors; an increase in the number of channels is good for robustness against noise, but comes at the high cost required to both prepare the physically necessary channels, but also to

biochemically activate those channels for transmission. My minimum and optimal encoding models significantly reduce this trade-off problem, by selecting channels essential for the current environment in a data-driven manner, as described below.

During navigation, individuals perform path planning, in which the question of where to move is more important than whether or not they can move. Thus, learning the locations of paths is more crucial than the locations of walls, to allow the participants to identify available directions. The minimum and optimal encoding models, representing participants' fMRI activities, had a tendency to assign active representation (i.e., '1') to scene views including paths (see Figs. 5.4 and 5.5). Moreover, the minimum and optimal encoding models allocated their code words so as to make the four scene views whose center square was a path more robust against bit inversion than the other four views whose center square was a wall (Figs. 5.2b, 5.4b and 5.5b). These results suggest that my minimum and optimal encoding models acquired a feature representation suitable for a subsequent decoding process by means of noisy information transmission; a similar encoding scheme could be implemented in the brain to produce stable decision making by decoding noisy information provided by its upstream channels. I also found that my data-driven minimum encoding model needed more channels than did a data-unrelated random encoding model that incorporated channels in a random order. This probably occurred because the GLM assigned larger beta weights to more important features. Such features, represented as encoding channels, would be placed with a high ranking according to my method. The more important features were likely correlated with each other, increasing redundancy in the list relative to that of the randomly ordered list. On the other hand, such redundancy in my data-driven encoders led to higher decoding performance versus that of the data-unrelated random encoders. This analysis suggests that my data-driven encoding models could detect important features, by paying for redundancy, and were effective for robust decoding against bit-inversion errors that can occur when decoding from noisy brain activities (Figs. 5.2c, 5.4c and 5.5c).

The frontal and/or parietal regions have been thought to be possible sources of anticipation in two contexts, reward prediction and future path planning (Ito et al., 2015; Chau et al., 2014; Noonan et al., 2010; Hunt et al., 2012). In the present study, I found that the GLM weights with the whole-scene model exhibited large values in five bilateral brain regions (Fig. 5.7a): rolandic operculum, SFG, IPG, and

precuneus, and TP. Moreover, the voxels showing a higher IGI were localized to the SFG and TP, and those showing higher prediction accuracy were localized to the IPG and precuneus (see Figs. 5.7b, 5.7c and 5.8). These results suggest that the fronto-parietal regions are involved in general aspects of scene anticipation, not just in reward prediction or path planning. Although the rolandic operculum was also identified in the GLM weight-based analysis, I found no further evidence that this region is related to encoding scene anticipation.

Although I did not directly investigate brain activities involved in path planning, scene anticipation should accompany path planning, especially in active navigation tasks such as MD. Actually, the decoding performance slightly increased over the progress of active navigation in MD (Fig. 5.9b), and scene anticipation was more accurate when a participant's performance was better. Moreover, when the participants followed a longer route than the optimal/suboptimal one, the predicted code words for scene views in the present route were similar to the correct code words of the scene views that they had experienced in the past (Figs. 5.11c and 5.12). These results imply that participants sometimes got lost, and could no longer predict upcoming scene views as far into the future. Furthermore, the observation that the future scene views could be predicted from the current brain activities (Fig. 5.11a, upper), especially when the participants were performing well in MD, may indicate that the participants' brain activities comprised prediction factors over multiple future steps. Such factors related to path planning might have lowered the decodability of one-step future scene views during active navigation in MD in comparison with passive navigation in SC (Fig. 5.9b).

Existing fMRI decoding studies presented encoding models embodied in higher visual areas (Stansbury, Naselaris, and Gallant, 2013; Nishimoto et al., 2011; Horikawa et al., 2013; Huth et al., 2012; Çukur et al., 2013). My present study focused on scene anticipation, which could be related to visual processing but is hypothetically a downstream visual process, and revealed that encoding can be localized in frontal and parietal areas that are known to be involved in decision making. We know very little of how information is represented for decision making, especially in higher brain regions. This lack of knowledge has motivated me to develop my data-driven ECOC encoders, the minimum and optimal encoding models, as particular implementations of a general error-tolerant encoding scheme. My analysis of robustness suggested that my data-driven encoding models could have detected important

features, thanks to their higher redundancy relative to the random channel list-based encoding models or the model with absolutely minimum number of channels (view-part model).

Previous studies have suggested that the human brain adopts parallel processing in the perceptual and behavioral stages, but serial processing in the decision stage (Bakkum et al., 2013; Pashler, 1994). The phenomenon of anticipation occurs midway between the perceptual stage and the step of behavior selection, and hence is considered to be in the early decision stage. The finding here that my optimal encoding model was more plausible than its deterministic counterparts, the whole-scene and view-part models, might suggest that scene anticipation is also processed in parallel so as to be error-tolerant. Notice, however, that fMRI has limited temporal resolution due to hemodynamic delay, so there remains the possibility that anticipation actually depends on serial processing. Meanwhile, the Bayesian brain hypothesis has been widely discussed in a number of studies (Knill and Pouget, 2004; Doya et al., 2006; Körding and Wolpert, 2006). This hypothesis assumes that the brain manipulates multiple possibilities by maintaining a probabilistic distribution (prior) of them. My current study is theoretically compatible with these studies.

My encoding models combined multiple binary channels linearly for representing fMRI activities. These constraints, linear combination and binarization, allowed me to describe my encoding models by ECOC. However, there may be more sophisticated encoding models in non-linear and/or non-binary domains. Seeking and describing such advanced encoding models would facilitate our understanding of the encoding mechanism of scene anticipation. Furthermore, scene anticipation in unfamiliar situations remains to be understood. The hippocampal place system contributes to the encoding of not only familiar situations, but also to a related novel experience occurring in the future (Dragoi and Tonegawa, 2011). Are such ‘preplay-like’ activities involved in scene anticipation mediated by the cerebral cortex? If yes, my decoding method may be able to detect individual expectations depending on his/her prior knowledge or character. Additionally, Although all my analyses in this study were all performed offline, the decoding analysis in the MD task can be extended to the online mode with the help of online fMRI measurement techniques; my encoder and decoder were solely determined by data from the SC task, which could be completed in advance of the MD task. Such online decoding technology may lead to the development of new brain-based devices to assist human navigation

and decision-making, for example, a secure guidance system that warns users of upcoming dangerous behaviors by notifying them of an incongruous belief that does not match reality.

Chapter 6

Conclusion and Future Directions

I used a non-invasive measurement modality, fMRI, and some multi-variate statistical analysis methods to deal with fMRI data. I here sum up this study by describing my contributions and remaining future directions toward further understanding of human anticipation. Main points of my study are as follows:

1. My computational model of navigational behaviors successfully described human decision making in an uncertain environment.
2. I successfully implemented the decoding of human anticipation even when the subjective belief is different from the objective reality.
3. Using the novel data-driven analysis, I suggested that the neural representation of anticipation is robustly and effectively encoded in the mPFC and parietal regions.

My methods have a potential to be widely applied to unveiling other higher-level cognitive functions in the brain.

6.1 Functional network for anticipation

Based on the findings of this study, I illustrate a possible functional network of anticipation (Fig. 6.1). Visual inputs (i.e., a current scene view) come from the OC. They are processed by well-known two major pathways; the dorsal stream and the ventral stream (Goodale and Milner, 1992). The spatial anticipation is mainly involved in the former stream, which is consistent with a widely accepted hypothesis that the dorsal stream would cause visually guided behaviors. The functional network is considered as consisting of two interactive connections. The mPFC includes both connections: the prefrontal-parietal network and the prefrontal-thalamo-hippocampal circuit.

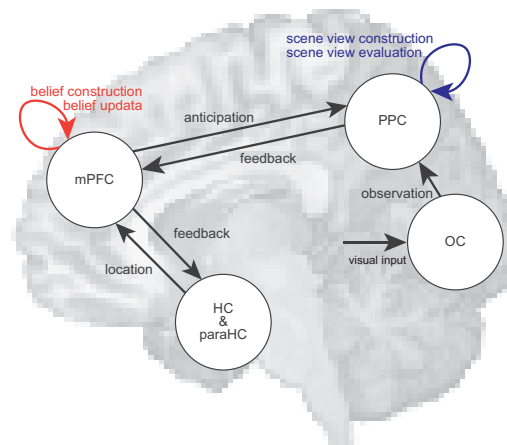


FIGURE 6.1: Diagram of anticipation network. The hippocampal place system represents a current position in the environment. This location information leads to a belief of upcoming observation in the mPFC. Based on the belief, the posterior parietal cortex constructs view anticipation as a visual image, and compares with a real scene view received through the OC.

Strong anatomical connections exist between the mPFC and the parietal regions. An interactive system of anticipation, belief construction and update, is implemented within this connectivity with cross-check of scene views. On the other hand, Ito et al., 2015 suggested that indirect projections from the mPFC are crucial for representation of the future path planning in rats' hippocampal place system. The mPFC sends information, via the nucleus reuniens, to the hippocampus. Thus, there could be an interactive network for performing a cross-check between the belief and the current observation.

6.2 Neural representation for anticipation

I proposed a new view of fMRI brain activities, in terms of noise-tolerant encoding. Using fMRI-data-driven analysis, I found reasonable encoding models that are biased towards environmental settings such as map topography. The brain encodes frequently-experience scene view with robust schemes. In contrast, indistinguishable representation is assigned to the scene views that are infrequent. This representation provides the noise-tolerant property, which allowed me to decode future scene views with higher decoding accuracies. This finding would suggest that our knowledge about information communication, such as the coding theory,

is useful for understanding brain functions that can be measured by non-invasive measurement modalities.

6.3 Challenges for the future

My contribution to anticipation raises as many questions as it has answered. I here list several remaining issues.

1. One of the most challenging remaining issues is causality of belief. In Chapter 4, I showed that the medial frontal and parietal regions reflect individual belief rather than objective reality. The next normative question is, then, how the hippocampal place system treats this incorrectness: where this inaccurate belief is generated.
2. In Chapter 5, I suggested that accurate behaviors could be associated with a long view toward the future. In our daily life, how we integrate various piece of information such as long-term prediction toward the future to produce best performance. Does an expert in the spatial navigation, like a taxi driver, has a longer-term prediction than an amateur, like a non-driver?
3. Toward applications, an on-line decoding system would be required. For many practical brain decoding systems, we need to apply a feedback to the brain in a closed-loop manner, such as, any real-world applications of brain-machine interface. On the other hand, the brain activities may change due not only to the feedback, but also to the feedback training. How the anticipation-related activities are modified through the closed-loop feedback is an important remaining issue.
4. A relationship between non-linearity is another challenging direction. I used an ECOC based encoding model, which implicitly assumed a linear encoding process consisting of multiple encoding pathways in the brain, to reproduce human fMRI activities. This is of course a simplification of the encoding process, but what would happen if we extend the encoding model into more general networks. Although the usage of such non-linear techniques is beyond the scope of my current study, due to the relatively short of human brain

activity data, I am also interested in the relationship human encoders and such non-linear machine learners.

Bibliography

- Albers, A. M. et al. (2013). "Shared representations for working memory and mental imagery in early visual cortex." In: *Current biology : CB* 23.15, pp. 1427–1431. ISSN: 1879-0445. DOI: [10.1016/j.cub.2013.05.065](https://doi.org/10.1016/j.cub.2013.05.065).
- Bakkum, D. J. et al. (2013). "Tracking axonal action potential propagation on a high-density microelectrode array across hundreds of sites." en. In: *Nature communications* 4:2181, doi: 10.1038/ncomms3181. ISSN: 2041-1723. DOI: [10.1038/ncomms3181](https://doi.org/10.1038/ncomms3181).
- Barto, A. G., Sutton, R. S., and Anderson, C. W. (1983). "Neuronlike adaptive elements that can solve difficult learning control problems". In: *IEEE Transactions on Systems, Man, and Cybernetics* SMC-13.5, pp. 834–846. ISSN: 0018-9472. DOI: [10.1109/TSMC.1983.6313077](https://doi.org/10.1109/TSMC.1983.6313077).
- Boland, P. J. (1989). "Majority systems and the condorcet jury theorem". In: *The Statistician* 38.3, p. 181. ISSN: 00390526. DOI: [10.2307/2348873](https://doi.org/10.2307/2348873).
- Bollinger, J. et al. (2010). "Expectation-driven changes in cortical functional connectivity influence working memory and long-term memory performance." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 30.43, pp. 14399–14410. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.1547-10.2010](https://doi.org/10.1523/JNEUROSCI.1547-10.2010).
- Bonnici, H. M. et al. (2012). "Detecting representations of recent and remote autobiographical memories in vmPFC and hippocampus." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 32.47, pp. 16982–16991. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.2475-12.2012](https://doi.org/10.1523/JNEUROSCI.2475-12.2012).
- Buzsáki, G. and Moser, E. I. (2013). "Memory, navigation and theta rhythm in the hippocampal-entorhinal system." In: *Nature neuroscience* 16.2, pp. 130–138. ISSN: 1546-1726. DOI: [10.1038/nn.3304](https://doi.org/10.1038/nn.3304).
- Chau, B. K. H. et al. (2014). "A neural mechanism underlying failure of optimal choice with multiple alternatives." In: *Nature neuroscience* 17.3, pp. 463–470. ISSN: 1546-1726. DOI: [10.1038/nn.3649](https://doi.org/10.1038/nn.3649).

- Courtney, S. M. et al. (1998). "An area specialized for spatial working memory in human frontal cortex". In: *Science* 279.5355, pp. 1347–1351. ISSN: 00368075. DOI: [10.1126/science.279.5355.1347](https://doi.org/10.1126/science.279.5355.1347).
- Çukur, T. et al. (2013). "Functional subdomains within human FFA." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 33.42, pp. 16748–66. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.1259-13.2013](https://doi.org/10.1523/JNEUROSCI.1259-13.2013).
- Daw, N. D., Niv, Y., and Dayan, P. (2005). "Uncertainty-based competition between prefrontal and dorsolateral striatal systems for behavioral control." In: *Nature neuroscience* 8.12, pp. 1704–11. ISSN: 1097-6256. DOI: [10.1038/nn1560](https://doi.org/10.1038/nn1560).
- Daw, N. D. et al. (2006). "Cortical substrates for exploratory decisions in humans." In: *Nature* 441.7095, pp. 876–9. ISSN: 1476-4687. DOI: [10.1038/nature04766](https://doi.org/10.1038/nature04766).
- Dickinson, A. (1985). "Actions and Habits: The Development of Behavioural Autonomy". In: *Philosophical Transactions of the Royal Society B: Biological Sciences* 308.1135, pp. 67–78. ISSN: 0962-8436. DOI: [10.1098/rstb.1985.0010](https://doi.org/10.1098/rstb.1985.0010).
- Dietterich, T. G. and Bakiri, G. (1995). "Solving multiclass learning problems via error-correcting output codes". In: *Journal of Artificial Intelligence Research* 2, pp. 263–286.
- Doll, B. B., Simon, D. A., and Daw, N. D. (2012). "The ubiquity of model-based reinforcement learning." In: *Current opinion in neurobiology* 22.6, pp. 1075–1081. ISSN: 1873-6882. DOI: [10.1016/j.conb.2012.08.003](https://doi.org/10.1016/j.conb.2012.08.003).
- Doll, B. B. et al. (2015). "Model-based choices involve prospective neural activity." In: *Nature neuroscience* 18.5, pp. 767–772. ISSN: 1546-1726. DOI: [10.1038/nn.3981](https://doi.org/10.1038/nn.3981).
- Doya, K. et al. (2006). *Bayesian Brain: Probabilistic Approaches to Neural Coding*. The MIT Press. ISBN: 026204238X.
- Dragoi, G. and Tonegawa, S. (2011). "Preplay of future place cell sequences by hippocampal cellular assemblies". In: *Nature* 469.7330, pp. 397–401. ISSN: 0028-0836. DOI: [10.1038/nature09633](https://doi.org/10.1038/nature09633).
- Epstein, R. A. (2008). "Parahippocampal and retrosplenial contributions to human spatial navigation." In: *Trends in cognitive sciences* 12.10, pp. 388–96. ISSN: 1364-6613. DOI: [10.1016/j.tics.2008.07.004](https://doi.org/10.1016/j.tics.2008.07.004).
- Eschenko, O. et al. (2008). "Sustained increase in hippocampal sharp-wave ripple activity during slow-wave sleep after learning". In: *Learning & Memory* 15.4, pp. 222–228. ISSN: 1072-0502. DOI: [10.1101/lm.726008](https://doi.org/10.1101/lm.726008).

- Ester, E. F., Serences, J. T., and Awh, E. (2009). "Spatially Global Representations in Human Primary Visual Cortex during Working Memory Maintenance". In: *Journal of Neuroscience* 29.48, pp. 15258–15265. ISSN: 0270-6474. DOI: [10.1523/JNEUROSCI.4388-09.2009](https://doi.org/10.1523/JNEUROSCI.4388-09.2009).
- Fiorillo, C. D. et al. (2003). "Discrete coding of reward probability and uncertainty by dopamine neurons." In: *Science (New York, N.Y.)* 299.5614, pp. 1898–902. ISSN: 1095-9203. DOI: [10.1126/science.1077349](https://doi.org/10.1126/science.1077349).
- Formisano, E. et al. (2008). "'Who' Is Saying 'What'? Brain-Based Decoding of Human Voice and Speech". In: *Science* 322.5903.
- Foster, D. J. and Wilson, M. A. (2006). "Reverse replay of behavioural sequences in hippocampal place cells during the awake state". In: *Nature* 440.7084, pp. 680–683. ISSN: 0028-0836. DOI: [10.1038/nature04587](https://doi.org/10.1038/nature04587).
- Fox, M. D. et al. (2005). "The human brain is intrinsically organized into dynamic, anticorrelated functional networks." In: *Proceedings of the National Academy of Sciences of the United States of America* 102.27, pp. 9673–9678. ISSN: 0027-8424. DOI: [10.1073/pnas.0504136102](https://doi.org/10.1073/pnas.0504136102).
- Fujisawa, S. et al. (2008). "Behavior-dependent short-term assembly dynamics in the medial prefrontal cortex." In: *Nature neuroscience* 11.7, pp. 823–833. ISSN: 1097-6256. DOI: [10.1038/nn.2134](https://doi.org/10.1038/nn.2134).
- Fyhn, M. et al. (2008). "Grid cells in mice". In: *Hippocampus* 18.12, pp. 1230–1238. ISSN: 10509631. DOI: [10.1002/hipo.20472](https://doi.org/10.1002/hipo.20472).
- Galati, G et al. (2000). "The neural basis of egocentric and allocentric coding of space in humans: a functional magnetic resonance study." In: *Experimental brain research* 133.2, pp. 156–164. ISSN: 0014-4819.
- Gervais, M. and Wilson, D. S. (2005). "The evolution and functions of laughter and humor: a synthetic approach." In: *The Quarterly review of biology* 80.4, pp. 395–430. ISSN: 0033-5770.
- Girardeau, G. et al. (2009). "Selective suppression of hippocampal ripples impairs spatial memory". In: *Nature Neuroscience* 12.10, pp. 1222–1223. ISSN: 1097-6256. DOI: [10.1038/nn.2384](https://doi.org/10.1038/nn.2384).
- Goodale, M. A. and Milner, A. (1992). "Separate visual pathways for perception and action". In: *Trends in Neurosciences* 15.1, pp. 20–25. ISSN: 01662236. DOI: [10.1016/0166-2236\(92\)90344-8](https://doi.org/10.1016/0166-2236(92)90344-8).

- Hafting, T. et al. (2005). "Microstructure of a spatial map in the entorhinal cortex". In: *Nature* 436.7052, pp. 801–806. ISSN: 0028-0836. DOI: [10.1038/nature03721](https://doi.org/10.1038/nature03721).
- Harris, K. D. et al. (2003). "Organization of cell assemblies in the hippocampus." In: *Nature* 424.6948, pp. 552–556. ISSN: 1476-4687. DOI: [10.1038/nature01834](https://doi.org/10.1038/nature01834).
- Harrison, S. A. and Tong, F. (2009). "Decoding reveals the contents of visual working memory in early visual areas." In: *Nature* 458.7238, pp. 632–635. ISSN: 1476-4687. DOI: [10.1038/nature07832](https://doi.org/10.1038/nature07832).
- Harvey, C. D., Coen, P., and Tank, D. W. (2012). "Choice-specific sequences in parietal cortex during a virtual-navigation decision task." In: *Nature* 484.7392, pp. 62–68. ISSN: 1476-4687. DOI: [10.1038/nature10918](https://doi.org/10.1038/nature10918).
- Hassabis, D. et al. (2009). "Decoding neuronal ensembles in the human hippocampus." In: *Current biology : CB* 19.7, pp. 546–554. ISSN: 1879-0445. DOI: [10.1016/j.cub.2009.02.033](https://doi.org/10.1016/j.cub.2009.02.033).
- Haxby, J. V. et al. (2001). "Distributed and overlapping representations of faces and objects in ventral temporal cortex." In: *Science (New York, N.Y.)* 293.5539, pp. 2425–30. ISSN: 0036-8075. DOI: [10.1126/science.1063736](https://doi.org/10.1126/science.1063736).
- Horikawa, T et al. (2013). "Neural decoding of visual imagery during sleep." In: *Science (New York, N.Y.)* 340.6132, pp. 639–642. ISSN: 1095-9203. DOI: [10.1126/science.1234330](https://doi.org/10.1126/science.1234330).
- Huettel, S. A., Song, A. W., and McCarthy, G. (2008). *Functional Magnetic Resonance Imaging*. 2nd editio. Sinauer Associates, p. 510. ISBN: 978-0878932863.
- Hunt, L. T. et al. (2012). "Mechanisms underlying cortical activity during value-guided choice." In: *Nature neuroscience* 15.3, pp. 470–476. ISSN: 1546-1726. DOI: [10.1038/nn.3017](https://doi.org/10.1038/nn.3017).
- Huth, A. G. et al. (2012). "A continuous semantic space describes the representation of thousands of object and action categories across the human brain." English. In: *Neuron* 76.6, pp. 1210–1224. ISSN: 1097-4199. DOI: [10.1016/j.neuron.2012.10.014](https://doi.org/10.1016/j.neuron.2012.10.014).
- Ito, H. T. et al. (2015). "A prefrontal–thalamo–hippocampal circuit for goal-directed spatial navigation". In: *Nature* 522.7554, pp. 50–55. ISSN: 0028-0836. DOI: [10.1038/nature14396](https://doi.org/10.1038/nature14396).
- Johnson, A. and Redish, A. D. (2007). "Neural ensembles in CA3 transiently encode paths forward of the animal at a decision point." In: *The Journal of neuroscience : the*

- official journal of the Society for Neuroscience* 27.45, pp. 12176–12189. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.3761-07.2007](https://doi.org/10.1523/JNEUROSCI.3761-07.2007).
- Kaelbling, L. P., Littman, M. L., and Cassandra, A. R. (1998). "Planning and acting in partially observable stochastic domains". In: *Artificial Intelligence* 101.1, pp. 99–134.
- Kamitani, Y. and Tong, F. (2005). "Decoding the visual and subjective contents of the human brain." In: *Nature neuroscience* 8.5, pp. 679–685. ISSN: 1097-6256. DOI: [10.1038/nn1444](https://doi.org/10.1038/nn1444).
- Kay, K. N. et al. (2008). "Identifying natural images from human brain activity." In: *Nature* 452.7185, pp. 352–355. ISSN: 1476-4687. DOI: [10.1038/nature06713](https://doi.org/10.1038/nature06713).
- Kitagawa, G. and Sato, S. (2001). "Monte Carlo Smoothing and Self-Organising State-Space Model". In: *Sequential Monte Carlo Methods in Practice*. New York, NY: Springer New York, pp. 177–195. DOI: [10.1007/978-1-4757-3437-9_9](https://doi.org/10.1007/978-1-4757-3437-9_9).
- Knill, D. C. and Pouget, A. (2004). "The Bayesian brain: the role of uncertainty in neural coding and computation." In: *Trends in neurosciences* 27.12, pp. 712–719. ISSN: 0166-2236. DOI: [10.1016/j.tins.2004.10.007](https://doi.org/10.1016/j.tins.2004.10.007).
- Körding, K. P. and Wolpert, D. M. (2006). "Bayesian decision theory in sensorimotor control." In: *Trends in cognitive sciences* 10.7, pp. 319–326. ISSN: 1364-6613. DOI: [10.1016/j.tics.2006.05.003](https://doi.org/10.1016/j.tics.2006.05.003).
- Marchette, S. A. et al. (2014). "Anchoring the neural compass: coding of local spatial reference frames in human medial parietal lobe." In: *Nature neuroscience* 17.11, pp. 1598–1606. ISSN: 1546-1726. DOI: [10.1038/nn.3834](https://doi.org/10.1038/nn.3834).
- Meyer, K. et al. (2010). "Predicting visual stimuli on the basis of activity in auditory cortices". In: *Nature Neuroscience* 13.6, pp. 667–668. ISSN: 1097-6256. DOI: [10.1038/nn.2533](https://doi.org/10.1038/nn.2533).
- Miyawaki, Y. et al. (2008). "Visual image reconstruction from human brain activity using a combination of multiscale local image decoders." In: *Neuron* 60.5, pp. 915–929. ISSN: 1097-4199. DOI: [10.1016/j.neuron.2008.11.004](https://doi.org/10.1016/j.neuron.2008.11.004).
- Naselaris, T. et al. (2009). "Bayesian reconstruction of natural images from human brain activity." In: *Neuron* 63.6, pp. 902–915. ISSN: 1097-4199. DOI: [10.1016/j.neuron.2009.09.006](https://doi.org/10.1016/j.neuron.2009.09.006).
- Naselaris, T. et al. (2011). "Encoding and decoding in fMRI." In: *NeuroImage* 56.2, pp. 400–410. ISSN: 1095-9572. DOI: [10.1016/j.neuroimage.2010.07.073](https://doi.org/10.1016/j.neuroimage.2010.07.073).

- Nishimoto, S. et al. (2011). "Reconstructing visual experiences from brain activity evoked by natural movies." In: *Current biology : CB* 21.19, pp. 1641–1646. ISSN: 1879-0445. DOI: [10.1016/j.cub.2011.08.031](https://doi.org/10.1016/j.cub.2011.08.031).
- Noonan, M. P. et al. (2010). "Separate value comparison and learning mechanisms in macaque medial and lateral orbitofrontal cortex." In: *Proceedings of the National Academy of Sciences of the United States of America* 107.47, pp. 20547–20552. ISSN: 1091-6490. DOI: [10.1073/pnas.1012246107](https://doi.org/10.1073/pnas.1012246107).
- Norman, K. A. et al. (2006). "Beyond mind-reading: multi-voxel pattern analysis of fMRI data." In: *Trends in cognitive sciences* 10.9, pp. 424–30. ISSN: 1364-6613. DOI: [10.1016/j.tics.2006.07.005](https://doi.org/10.1016/j.tics.2006.07.005).
- O'Keefe, J. and Dostrovsky, J. (1971). "The hippocampus as a spatial map. Preliminary evidence from unit activity in the freely-moving rat". In: *Brain Research* 34.1, pp. 171–175. ISSN: 00068993. DOI: [10.1016/0006-8993\(71\)90358-1](https://doi.org/10.1016/0006-8993(71)90358-1).
- Owen, A. M. (1997). "Cognitive planning in humans: Neuropsychological, neuroanatomical and neuropharmacological perspectives". In: *Progress in Neurobiology* 53.4, pp. 431–450. ISSN: 03010082. DOI: [10.1016/S0301-0082\(97\)00042-7](https://doi.org/10.1016/S0301-0082(97)00042-7).
- Pashler, H. (1994). "Dual-task interference in simple tasks: Data and theory." In: *Psychological Bulletin* 116.2, pp. 220–244.
- Pearce, J. M. and Hall, G. (1980). "A model for Pavlovian learning: Variations in the effectiveness of conditioned but not of unconditioned stimuli." In: *Psychological Review* 87.6, pp. 532–552. ISSN: 0033-295X. DOI: [10.1037/0033-295X.87.6.532](https://doi.org/10.1037/0033-295X.87.6.532).
- Penner, M. R. and Mizumori, S. J. (2012). "Neural systems analysis of decision making during goal-directed navigation". In: *Progress in Neurobiology* 96.1, pp. 96–135. ISSN: 03010082. DOI: [10.1016/j.pneurobio.2011.08.010](https://doi.org/10.1016/j.pneurobio.2011.08.010).
- Pfeiffer, B. E. and Foster, D. J. (2013). "Hippocampal place-cell sequences depict future paths to remembered goals." In: *Nature* 497.7447, pp. 74–79. ISSN: 1476-4687. DOI: [10.1038/nature12112](https://doi.org/10.1038/nature12112).
- Quinlan, J. R. (1986). "Induction of decision trees". In: *Machine Learning* 1.1, pp. 81–106. ISSN: 0885-6125. DOI: [10.1007/BF00116251](https://doi.org/10.1007/BF00116251).
- Reddy, L., Tsuchiya, N., and Serre, T. (2010). "Reading the mind's eye: decoding category information during mental imagery." In: *NeuroImage* 50.2, pp. 818–825. ISSN: 1095-9572. DOI: [10.1016/j.neuroimage.2009.11.084](https://doi.org/10.1016/j.neuroimage.2009.11.084).

- Rescorla, R. A. and Wagner, A. R. A. (1972). "A theory of Pavlovian conditioning: The effectiveness of reinforcement and non-reinforcement". In: *Classical conditioning II: current research and theory*. Ed. by A. H. Black and W. F. Prokasy. Appleton-Century-Crofts, New York. Chap. 3, pp. 64–99.
- Rodriguez, P. F. (2010). "Neural decoding of goal locations in spatial navigation in humans with fMRI." In: *Human brain mapping* 31.3, pp. 391–397. ISSN: 1097-0193. DOI: [10.1002/hbm.20873](https://doi.org/10.1002/hbm.20873).
- Schultz, W (1986). "Responses of midbrain dopamine neurons to behavioral trigger stimuli in the monkey." In: *Journal of neurophysiology* 56.5, pp. 1439–61. ISSN: 0022-3077.
- Schultz, W., Apicella, P, and Ljungberg, T (1993). "Responses of monkey dopamine neurons to reward and conditioned stimuli during successive steps of learning a delayed response task". In: *J. Neurosci.* 13.3, pp. 900–913.
- Schultz, W., Dayan, P, and Montague, P. R. (1997). "A neural substrate of prediction and reward." In: *Science (New York, N.Y.)* 275.5306, pp. 1593–1599. ISSN: 0036-8075.
- Sigman, M. and Dehaene, S. (2008). "Brain mechanisms of serial and parallel processing during dual-task performance." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 28.30, pp. 7585–7598. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.0948-08.2008](https://doi.org/10.1523/JNEUROSCI.0948-08.2008).
- Simon, S. R. et al. (2002). "Spatial Attention and Memory Versus Motor Preparation: Premotor Cortex Involvement as Revealed by fMRI". In: *J Neurophysiol* 88.4, pp. 2047–2057.
- Skinner, B. F. (1938). *The behavior of organisms: an experimental analysis*. Appleton-Century-Crofts, New York.
- Stansbury, D. E., Naselaris, T., and Gallant, J. L. (2013). "Natural scene statistics account for the representation of scene categories in human visual cortex." In: *Neuron* 79.5, pp. 1025–1034. ISSN: 1097-4199. DOI: [10.1016/j.neuron.2013.06.034](https://doi.org/10.1016/j.neuron.2013.06.034).
- Sutton, R. S. (1988). "Learning to predict by the methods of temporal differences". In: *Machine Learning* 3.1, pp. 9–44. ISSN: 0885-6125. DOI: [10.1007/BF00115009](https://doi.org/10.1007/BF00115009).
- Takenouchi, T. and Ishii, S. (2009). "A multiclass classification method based on decoding of binary classifiers." In: *Neural computation* 21.7, pp. 2049–2081. ISSN: 0899-7667. DOI: [10.1162/neco.2009.03-08-740](https://doi.org/10.1162/neco.2009.03-08-740).

- Tobler, P. N., Fiorillo, C. D., and Schultz, W. (2005). "Adaptive coding of reward value by dopamine neurons." In: *Science (New York, N.Y.)* 307.5715, pp. 1642–5. ISSN: 1095-9203. DOI: [10.1126/science.1105370](https://doi.org/10.1126/science.1105370).
- Tong, F. et al. (2012). "Relationship between BOLD amplitude and pattern classification of orientation-selective activity in the human visual cortex." In: *NeuroImage* 63.3, pp. 1212–1222. ISSN: 1095-9572. DOI: [10.1016/j.neuroimage.2012.08.005](https://doi.org/10.1016/j.neuroimage.2012.08.005).
- Tzourio-Mazoyer, N et al. (2002). "Automated anatomical labeling of activations in SPM using a macroscopic anatomical parcellation of the MNI MRI single-subject brain." In: *NeuroImage* 15.1, pp. 273–289. ISSN: 1053-8119. DOI: [10.1006/nimg.2001.0978](https://doi.org/10.1006/nimg.2001.0978).
- Vallar, G et al. (1999). "A fronto-parietal system for computing the egocentric spatial frame of reference in humans." In: *Experimental brain research* 124.3, pp. 281–286. ISSN: 0014-4819.
- Vass, L. K. and Epstein, R. A. (2013). "Abstract representations of location and facing direction in the human brain." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 33.14, pp. 6133–6142. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.3873-12.2013](https://doi.org/10.1523/JNEUROSCI.3873-12.2013).
- Vintch, B. and Gardner, J. L. (2014). "Cortical correlates of human motion perception biases." In: *The Journal of neuroscience : the official journal of the Society for Neuroscience* 34.7, pp. 2592–2604. ISSN: 1529-2401. DOI: [10.1523/JNEUROSCI.2809-13.2014](https://doi.org/10.1523/JNEUROSCI.2809-13.2014).
- Yamashita, O. et al. (2008). "Sparse estimation automatically selects voxels relevant for the decoding of fMRI activity patterns." In: *NeuroImage* 42.4, pp. 1414–1429. ISSN: 1095-9572. DOI: [10.1016/j.neuroimage.2008.05.050](https://doi.org/10.1016/j.neuroimage.2008.05.050).
- Yoshida, W. and Ishii, S. (2006). "Resolution of uncertainty in prefrontal cortex." In: *Neuron* 50.5, pp. 781–789. ISSN: 0896-6273. DOI: [10.1016/j.neuron.2006.05.006](https://doi.org/10.1016/j.neuron.2006.05.006).
- Yukinawa, N. et al. (2009). "Optimal aggregation of binary classifiers for multiclass cancer diagnosis using gene expression profiles." In: *IEEE/ACM transactions on computational biology and bioinformatics / IEEE, ACM* 6.2, pp. 333–343. ISSN: 1557-9964. DOI: [10.1109/TCBB.2007.70239](https://doi.org/10.1109/TCBB.2007.70239).

- Yushkevich, P. A. et al. (2006). "User-guided 3D active contour segmentation of anatomical structures: significantly improved efficiency and reliability." In: *NeuroImage* 31.3, pp. 1116–1128. ISSN: 1053-8119. DOI: [10.1016/j.neuroimage.2006.01.015](https://doi.org/10.1016/j.neuroimage.2006.01.015).