

声における重ね合わせに関する認知科学研究の可能性： 量子論的世界観を踏まえた自己観と人間観の展望

林 大輔^{1,*}

¹ 日本たばこ産業株式会社

Cognitive science research about superposition in voice:
Possible future views of self and humanity based on quantum perspective

Daisuke Hayashi^{1,*}

¹ Japanese Tobacco Inc.

This paper explores the potential of cognitive science research into "superposition in voice," a concept inspired by quantum theory, to examine how voices can simultaneously embody multiple identities or interpretations. Focusing on three perspectives (self, individuality, and humanity), the paper refers to various research relating to each perspective in multiple disciplines. It shows that the advancements in human skills related to voice in pop culture such as voice actor and VTuber, as well as in engineering technologies such as speech synthesis and voice conversion, are blurring the boundaries between "self and others," "speaker and character," and "human and machine." It suggests that these blurred boundaries can be viewed as "superposition in voice." Based on this point of view, the paper discusses the various questions relating to each of the three perspectives, and proposes the possible research approaches in cognitive science, such as quantum cognition and projection science. It also highlights the importance of a digital archive of voices in pop culture, as well as the ethical, legal, and social issues surrounding voice engineering. Overall, the paper advocates interdisciplinary research that bridges the gap between the humanities and engineering in order to redefine future views of the self and humanity.

Keywords: voice (声), superposition (重ね合わせ), pop culture (ポップカルチャー), voice engineering technology (音声工学技術), view of self (自己観), view of humanity (人間観)

Received xx June 2025; Accepted xx xxx 2025; Available online at J-STAGE xx xxx 2026

1. はじめに

声は日常にあふれており、私たちは当たり前のように声に触れ、声を用いて日々を過ごしている。私たちは、「私の声」というものを持っていると感じ、声を聞いた時にそれが「あの人の声である」と認識できていると思い、声でコミュニケーションしている相手が「誰であるのか」を理解しているように思われる。しかし、声は曖昧性を内包しており、自己と他者、話者とキャラクター、あるいは人間と機械といった、時には一見相反するもの同士の“間”をゆらぎうる。それらのゆらぎは、ポップカルチャーにおける人間の声の技術、あるいは音声合成や声質変換などの音声

工学技術によって様々な形で生じている。そのようなゆらぎの状態を、量子論的世界観に基づいて、複数の解釈の「重ね合わせ」状態として考えることで、量子的な状態の「重ね合わせ」と関連付けて検討でき、認知科学の観点からそのような状態を表現し、説明するような研究を行える可能性がある（布山・西郷, 2022）。

本稿では、声における自己と他者の“間”と関わる「私らしさ」、話者とキャラクターの“間”と関わる「その人らしさ」、そして人間と機械の“間”と関わる「人間らしさ」という3つの観点について、声を題材にこれまでに行われてきた研究を概観しつつ、そ

* 責任著者。 E-mail: daisuke.hayashi@jt.com

1 れぞれ声における重ね合わせがどのように生じてい
2 るのか、その点についてどのような論考があるのか
3 について検討を行う。その上で、今後どのような研究
4 の方向性が考えられるのかについて展望を述べ、量子
5 論的世界観に基づいた声に関する学際的な研究によ
6 り、新たな「自己観」や「人間観」がどのように提起
7 されうるのかについて議論を行う。

8 2. 声における「私らしさ」

9 私らしさ、すなわち自己に関する認知は大きな学
10 問上の問いであり、これまでに様々な研究が行われ
11 てきている（嶋田, 2019; 田中, 2023）。私たちは声を
12 発する時、自分の声を常に聞いているため、声は自己
13 を規定する重要な 1 つの要因たりうる。本章では、
14 声と自己との関係についての研究を概観した上で、
15 声において自己と他者がいかに重なりうるのかにつ
16 いて検討を行う。

17 2.1 「私らしさ」の認知

18 声は呼吸が声帯を振動させ、その振動が声道を通
19 って口から外に出て、空気を振動させながら伝わり、
20 聞き手の耳に届いて「声」として知覚されるものであ
21 る。そのように、口から出て空気を伝わって耳に入っ
22 てくる音を気導音と呼ぶが、声を発した本人である
23 話者が聞くのは気導音だけではない。発声時の声帯
24 振動が頭蓋骨を通じて、直接内耳に伝わってきた音
25 も話者は聞いており、そのような音は骨導音と呼ば
26 れる（日本音響学会, 1996）。そして気導音と骨導音が
27 合わさった「声」が、普段話者自身が聞いている声（自
28 己聴取音声）であり、これが通常、私たちにとっての
29 「自分の声」である。

30 骨導音があることで、自己聴取音声は話者以外の
31 聞き手が聞いている声とは異なっており、録音した
32 自身の声を聞くといつもと違って聞こえるため「自
33 分の声じゃないように感じる」ことになる。そのよう
34 ないつもと違う録音音声について、日常的な感覚と
35 して「好きではない」という話を聞くことは多く、実
36 際、先行研究では自身の録音音声に対してネガティ
37 ブな反応を示すことが報告されている（Daryadar &
38 Raghibi, 2015; Holzman & Rousey, 1966）。一方で、自
39 分の声は他者の声に比べて高く評価するという研究
40 も存在している（Hughes & Harrison, 2013）。また、自
41 分の声だけでなく、自分に似た声も評価が高いとも
42 報告されており（Peng et al., 2019）、自分の声に似た
43 アバターから教わる方が、メタバース上での教育ゲ

44 ームの効果が高いという研究も存在する（Kao et al.,
45 2021）。このように話者自身の録音音声に対する知見
46 は多岐にわたっており、個人特性と録音音声に対す
47 る評価との関係についても検討が行われている
48 （Lundh et al., 2002; Yanagida et al., 2023）。

49 録音音声ではなく、骨導音込みの自己聴取音声に
50 着目した研究も行われている。自己聴取音声は、気導
51 音に比べて低周波数ほど強く、高周波数ほど弱く聞
52 こえると言われているが（日本音響学会, 1996）、様々
53 なフィルタを用いて録音音声を変換し、自己聴取音
54 声との類似度を評価した研究では、どのフィルタが
55 最も自己聴取音声と類似した音声を生成できるかに
56 ついては個人差が大きいことが報告されている
57 （Kimura & Yotsumoto, 2018）。また関連して、録音音
58 声を複数のフィルタで変換した音声を用いて、自己
59 聴取音声と最も類似した音声を聞いている時と最も
60 異なる音声を聞いている時の脳活動を fMRI を用い
61 て比較し、「自分の声」の処理と関連する脳領域を調
62 べた研究も行われている（Hosaka et al., 2021）。近年
63 では、自己聴取音声を工学的に生成することで、「自
64 分の声」に講義をしてもらった際の課題成績を調べ
65 た研究や（福田他, 2025）、「自分の声」を用いて英語
66 の発音学習を行った研究など（山中他, 2025）、自己聴
67 取音声の工学応用に関する研究も行われている。

68 では、そのような声における「私らしさ」は変容し
69 うのだろうか。この点について、Gallagher (2000)
70 におけるミニマル・セルフとナラティブ・セルフそれ
71 ぞれと関連すると思われる研究が存在している。ミ
72 ニマル・セルフは、記憶などを失ってもなお残る「私」
73 という感覚であり、行為主体感と身体所有感からな
74 る。このうち、「この声は私が発したものである」と
75 という感覚である行為主体感と関連して、声の高さを
76 工学的に操作してフィードバックすると、声に対す
77 る行為主体感が弱くなることが報告されている
78 （Ohata et al., 2022）。また、本来は自分で発した声で
79 あるにも関わらず、声に対する行為主体感が失われ
80 ることで、発話主が誤帰属されて他者の声として感
81 じられ、それが幻聴として報告されるという幻聴の
82 自己モニタリング仮説も提案されている（浅井・丹野,
83 2010）。さらに昨今では、「自身の声を他者の声に変換
84 してリアルタイムでフィードバックする」声質変換
85 技術が発展しており（Sisman et al., 2021）、発話に対
86 して他者の声をフィードバックすると、他者の声を
87 自分の声のように感じる、すなわち他者の声に対す

1 る所有感が生じる可能性が示唆されている（中川・田
2 中, 2021）。そのような“錯覚”を活用して、フィード
3 バック音声をよく知られているキャラクターの声に
4 変換した際の認知課題の成績を検討した研究や
5 （Kunimi et al., 2024）、遠隔操作でロボット接客をす
6 る際に、ロボットらしい声をフィードバックするこ
7 とで、「ロボットになりきりやすくなる」ことを示し
8 た研究などが存在している（Ogawa et al., 2024）。

9 これらの研究では、変換した音声のリアルタイム
10 フィードバックを行うことで、行為主体感を含めた
11 ミニマル・セルフに働きかけつつ、「私らしさ」とい
12 う感覚をも変容させている、すなわちなラティブ・セ
13 ルフにも影響している可能性が考えられる。ナラテ
14 イブ・セルフは、アイデンティティとしての自己でも
15 あり、自伝的記憶などを含めた「私らしさ」という感
16 覚である。より長期的に声を変換して日常を過ごし
17 ている方々を対象とした研究として、國見他（2024）
18 では、ソーシャル VR プラットフォームの 1 つであ
19 る VRChat において、声質変換技術あるいは発声の工
20 夫を用いて「身体的な性別とは異なる性別の声」を日
21 常的に発しているユーザを対象にアンケート調査を
22 行っている。その結果、声を変えて話した際に「別人
23 になったような感覚」があること、一方でそのような
24 場合でも「現実世界での自分と繋がりがあある」とは感
25 じることが示された。そのように、ソーシャル VR 上
26 で理想のかわいさを表現するために、声もかわいい
27 ものとなるように努力している方々が存在しており
28 （Bredikhina, 2022）、長期的に声を変えて過ごすこと
29 で「私らしさ」という感覚も変容しうることが示唆さ
30 れている。そのようなユーザの 1 人であるバーチャ
31 ル美少女ねむは、「自分の声が変わった瞬間、ついさ
32 っきまで普通に仕事をしていたはずの自分の心がく
33 るりと裏返って、美少女としての私が目覚めます…
34 ボイチェンは、私の心の美少女スイッチなのです（バ
35 ーチャル美少女ねむ, 2022, p.179）」と述べており、声
36 を変えることは、ミニマル・セルフだけでなくナラテ
37 イブ・セルフにも影響しうることが示唆されている。

38 2.2. 自己と他者の重ね合わせ

39 前節において概観した、声における「私らしさ」の
40 認知に関する研究を踏まえつつ、本節では、声におい
41 て「自己」と「他者」の重ね合わせが生じている可能
42 性について検討を行う。

43 前節のおわりに述べたように、ソーシャル VR 上
44 で声を変えて話しているユーザは、「別人になったよ

45 うな感覚」を抱きつつ、同時に現実の自己との連続性
46 は感じている（國見他, 2024）。「現実と VRChat の自
47 己は別々の自己であると認識しながらも自己連続性
48 を感じている可能性がある（國見他, 2024, p. 4）」こと
49 は、「自己でもあるが、他者でもある」状態、すなわ
50 ち自己と他者の重ね合わせ状態が生じていることを
51 示唆している。また、國見他（2024）では身体的な性
52 別とは異なる性別の声を用いているユーザを対象と
53 していたが、ソーシャル VR のユーザの中には、見た
54 目（アバター）はかわいい美少女でありつつ、声は地
55 声のままである男性ユーザも存在している
56 （Bredikhina, 2022）。これは見た目と声が不一致な状
57 態であり、違和感を抱きうるとも考えられるが、「現
58 在のソーシャル VR では『女性アバターから男性の
59 声が出る』ようなパターンが非常に多く、それもだん
60 だんと当たり前のこととして受け入れられるように
61 なってきています（バーチャル美少女ねむ, 2022,
62 p.181）」との言及もあり、見た目の性別と声の性別が
63 異なる状態は「女性であり、男性でもある」という重
64 ね合わせ状態として、そのまま受け入れられている
65 可能性が考えられる。

66 ソーシャル VR 上でのこれらの重ね合わせは、技
67 術やサービスの発展に伴って昨今生じているものだ
68 と考えられるが、以前から、「声を用いて別人になる」
69 ことは行われており、その顕著な例がアニメにおけ
70 る声優だと考えられる。声優は声でキャラクターを
71 演じ、表現するが、たとえば千葉繁は「役に自分が寄
72 るシステムと、自分に役を引きつけるシステム、ふた
73 つの演技システムがあるんです。僕はいつも後者で
74 演技をやっています。こうなると見た目のキャラク
75 ターが違うだけで、どの役も本質的には自分なんで
76 す（藤津, 2019, p.59）」と述べており、キャラクター
77 という別人になりつつも、それも自分である、すなわ
78 ち「キャラクターであるが、自己でもある」ような重
79 ね合わせ状態だと捉えることができる。また、そのよ
80 うに「声で演じる」ことは声優に限った話ではなく、
81 「演劇や歌では、声はキャラクターやパーソナリテ
82 ィーを演じるための表現の一部として、意識的に用
83 いられる…声は仮面や化粧と同じように、素顔を隠
84 し、自分とは違った別の時代、別の身分、別の性格を
85 もつ役柄を演じるための道具の一つである（増野,
86 2014, p. 196)」。そして演じる行為を行うことは、自己
87 ではない何者かになるというよりは、「自己の二重化、
88 つまり『私でありながら、私でないもの（たとえば役

柄)でもある』状態(増野, 2014, p.205)」であると述べられており, 「自己であり, 他者でもある」状態を生じさせることこそが, 演じることだとも捉えることができる。認知神経科学研究において, 物語の登場人物になり切っている時には自分自身について考えるのと同じ部位が活動することが示されており(Broom et al., 2021), 一方で“fictional”な一人称視点にいる時, つまり何かの役になっている時には, 全体として脳活動が抑えられ, 自己という感覚が弱まる可能性を示唆する知見も存在することから(Brown et al., 2019), 演じている時の自己は, 自己でありつつ自己ではない状態だと考えることができそうである。

そして声で演じることが, 実は私たちにとって特別なものではなく, 置かれている場面や自身の役割に合わせて意識的あるいは無意識的に声質を操作することで, 望ましい印象を与えることは日常的に行われている(モクタリ, 2008)。たとえば, 電話で話す時と対面で話す時で声質が違っていたり, 初対面の相手と話す時と親しい相手と話す時で声質が変わったりすることは珍しくない。また, 百貨店の売り子, 電車の車内アナウンスなど, 職業に応じてそれに合う声質, 発声が存在しているようにも見受けられる(Grawe, 2018)。このように私たちは, 仕事などの社会的役割に応じて発せられる「自発型演技音声」を使っており(宮島他, 2013), 「たとえば商売人の声も, その人自身の声でありながら, その人のふだんの生活の声とは別のものでもある。限られた文脈の中で, 自分の声を変形させることで, 商売や職業に適したパーソナリティを確立する(増野, 2014, p.205)」との言及にもある通り, 私たちは声を“変える”ことで, 自己と「演じるべき役」を日常的に重ね合わせているのかもしれない。

3. 声における「その人らしさ」

私たちは声を聞いた時に, 話者の年齢や性別, 性格印象などの「どのような人なのか」に関する情報(話者情報)を知ることができる。また, 個人の同定と関連するような, より個々の話者に特有の「誰なのか」に関する情報(個人性情報)についても知ることができる。それらの情報と, 声道などの生理学的特徴, その特徴に起因する音声の音響特徴, そしてヒトの知覚的な特性との関係について, 様々な研究が行われてきている(本多, 2018; 森他, 2014; Schweinberger et al., 2014)。そしてこれらの話者情報と個人性情報を統合して得られる認知が, 声における「その人らしさ」

の認知だと考えられる。本章では, 声とその人らしさとの関係についての研究を概観した上で, 声において話者とキャラクターがいかに重なりうるのかについて検討を行う。

3.1 「その人らしさ」の認知

話者情報に含まれる年齢や性別については, 加齢に伴う成長や性差に基づく体格差に依存して, 声帯や声道などの調音器官に変化が生じ, そのことが音響特徴に, ひいては知覚に影響する。たとえば, 子供よりも大人の方が, あるいは女性よりも男性の方が, 平均的に体が大きいため, 声帯も声道も長く, 結果として基本周波数とスペクトル重心が低くなる(栗田, 1988; 森他, 2014; Titze, 1989; 内田・森勢, 2020; Vorperian et al., 2009; Vorperian et al., 2011)。生理学的特徴や音響特徴だけでなく, 音声聞いた際, 話者の年齢や体格などの情報がある程度正確に知覚できることを示した研究や(Krauss et al., 2002), 声の性別に対する順応と残効が生じることを示した研究が存在しており(Schweinberger et al., 2008), 話者情報が認知的に処理されていることが示されている。

声を聞いた際に話者に対して抱く印象も話者情報の一部だと考えると, 広義の声質(音声波から知覚される韻質以外の聴覚上の特質(森他, 2014; 日本音響学会, 2003, p. 194))は話者の性格印象に影響を与えることが, 様々な研究において示されており(Aronovitch, 1976; Belin et al., 2017; McAleer et al., 2014; Scherer, 1972, 1978; 内田, 2000, 2002, 2005a, 2005b, 2006, 2009, 2011; 内田・中畝, 2004), 声に基づいた人物像の知覚の3次元モデルなども提案されている(勅使河原, 2019)。また, 声質の印象についても, 日常表現語を用いた記述が行われている(木戸・粕谷, 1999)。

次に, 名状しがたい「この人の声」という感覚と関わる個人性情報については, 生理学的特徴の個人差に着目した検討が行われている。たとえば声道の形状の個人差と関連する音響特徴を検討し, 高周波成分における違いが個人の同定の手がかりとなることが示されている(本多, 2018; Kitamura et al., 2005)。一方で, 話者によって個人の同定に関わる手がかりが異なる可能性も指摘されている(Van Lancker, Kreiman, & Emmorey, 1985)。認知的な処理としては, 声の個人性に対する順応と残効が生じることを示した研究が存在している(Latinus & Belin, 2011)。

これらの話者情報と個人性情報は完全に独立な関

1 係ではない。たとえば、記憶課題においてターゲット
2 音声と似た異なる音声を「ターゲット音声である」と誤
3 認して回答することが増えてしまう(井上, 2022)。ある
4 いは、年齢や性別などの話者情報が違って聞こえる
5 複数の音声の中から同一の話者の音声を正しく弁別
6 することは、話者情報の違いを弁別するよりも難し
7 いことが示唆されている(林, 2019)。これらの研究は、
8 話者情報を無視して個人性情報のみに基づいて「誰
9 の声か」を認識することの難しさを示しており、話者
10 情報と個人性情報が合わさって「その人らしさ」の認
11 知が形作られていることの傍証だと考えられる。

12 そしてアニメを始めとしたポップカルチャーの中
13 では、話者情報を人間の声の技術によって適切に操
14 作することで、キャラクター表現がなされている。そ
15 のようなキャラクター演技音声を対象とした研究と
16 しては、同一声優による異なる性格を持つキャラク
17 ターの演技音声を分析した研究や(佐藤, 2023)、同一
18 声優の異なる年齢や性別を演じた音声を比較した研
19 究(林, 2019; 林・森勢, 2025; 丸島, 2020a, 2020b)、キ
20 ャラクターの声のステレオタイプ性と音響特徴量と
21 の関係について調べた研究(石井・伊藤, 2019)、善玉
22 と悪玉の声質の違いを調べた研究(Teshigawara, 2004;
23 勅使河原他, 2005)、実際のアニメ映画における異な
24 る性格のキャラクターの声を比較した研究
25 (Crochiquia et al., 2020)などが存在している。ある
26 いは、メイド喫茶の「ツン」と「萌え」という異なる
27 タイプのメイドそれぞれがどのような発話をしてい
28 るかを調べる研究も行われている(Kawahara, 2016)。
29 このような人間の声の技術による話者情報の操作は、
30 ポップカルチャーにおけるキャラクター表現の重要
31 な要素だと考えられる。

32 そして話者情報を操作してキャラクターを表現し
33 ても、話者が同一の個人であれば、個人性情報、
34 すなわち「誰であるのか」は同一であり、その情報は
35 声に埋め込まれていると考えられる。すなわち、「ア
36 ニメーションにおいてキャラクターが発する音声は、
37 現実に存在する声優によって発せられた音声でもあ
38 る(稲岡, 2016, p. 199)」のだが、「キャラクターの声
39 は『役者の声』として視聴者に届いてはいない…役者
40 の身体から発しつつも、『図像』の力によって、役者
41 の身体から切り離されて視聴者に響いているのがキ
42 ャラクターの声である(藤津, 2018, pp. 112-113)」。

43 優は、「画面内のキャラクターの動きや演出に合わせ
44 て、声の高低や濃淡などを変えたりする技術や、感情
45 を音声に乗せる技術など、細かな基礎的技術の集積
46 (稲岡, 2016, p. 120)」により、キャラクターの音声は
47 現実に存在する声優によって発せられた音声である
48 ことを「鑑賞者に意識させずにキャラクターによっ
49 て発せられた音声として鑑賞者に認識させるという
50 特質(稲岡, 2016, p. 199)」を持っている。このこと
51 は、声優自身の語りとしても表れており、たとえば声
52 優の久川綾は「このキャラを久川綾がやっているん
53 だと意外に思われたい(藤津, 2017, p. 238)」と述べ、
54 また声優の森川智之も「僕の声優観というのは、自分
55 があまり出てはいけないというものです…僕のファ
56 ンの方が僕目的で映画館に入ったとしても、それを
57 すっかり忘れて作品を楽しめるように演技したいと思
58 っています(藤津, 2019, pp. 177-178)」と述べている。
59 このように、本来は個人性情報が同一である中で、そ
60 れを意識させない、すなわち身体を伴った人間の声
61 ではなくキャラクターの声として聞こえるように話
62 者情報を適切に操作することが、ポップカルチャー
63 における人間の声の技術だと考えられる。

64 3.2 話者とキャラクターの重ね合わせ

65 前節において概観した、声における「その人らしさ」
66 の認知に関する研究を踏まえつつ、本節では、特にポ
67 ップカルチャーにおける人間の声の技術に着目して、
68 声において「話者」と「キャラクター」の重ね合わせ
69 が聞き手の中で生じている可能性について、複数の
70 状況に分けて検討を行う。

71 3.2.1 アニメにおける声優の声

72 まず始めに、前節でも言及した、アニメにおける
73 声優の声に着目する。前節で記載した通り、声優のキ
74 ャラクター演技音声は、声優自身の声でありつつ、そ
75 れを意識させず、キャラクターの声として聞こえるよ
76 うに技術が用いられている。しかし、「アニメーショ
77 ンにおいてキャラクターによって発せられたとみな
78 すことのできる音声は、現実に存在する声優と作品
79 内に存在するキャラクターという二重の発生源を有
80 する(稲岡, 2016, p. 122)」以上、キャラクターの声と
81 してしか聞こえないとは言えず、そこに身体を持っ
82 た人間としての声優を感じることは避けがたい。「現
83 実に私たちがアニメーションを鑑賞する仕方を思い
84 起こすと、『声優の声は声優の声であることを隠す』
85 という点が常に重視されているとは必ずしも断言で
86 きない…作品を鑑賞する際にも、画面に連動して聴

こえてくる音声がどの声優によって発せられたものかを十分に承知した上で鑑賞するという場合も珍しくはない(稲岡, 2016, p. 121)」のが実際であり、「聴取者たちは、一方においてはキャラクターの同一性と物語世界を成立させる声として声優の声を聞いており、他方においては、声優が各作品を横断して複数のキャラクターと世界を成立させている単一の声であると了解している(黒寄, 2018, p. 190)」。アイコン・インデックスという話者情報・個人性情報とそれぞれ関連が深い概念を導入した論考(ジン, 2010)を踏まえて、「視聴者は特定の声を通じてキャラクターの『人格・のようなもの』を了解すると同時に、声の『インデックス』性に基づく声優の身体を思い浮かべる(鈴木, 2014, p. 112)」とも述べられている。そしてそれは、「アニメーションにおいてキャラクターが発する音声を、キャラクターが発したものと受け取るか、声優が発したものと受け取るか、概念上は区別できても、それを聴覚的印象の水準で区別することはきわめて困難である(稲岡, 2016, p. 123)」ことに起因しているとも考えられ、その結果として「テレビアニメを観るとき…実際の発話者から声が発せられているとも、キャラクターから声が発せられているとも、どちらにも確信を得ていない(黒寄, 2018, p. 194)」い状態が生まれる。すなわち、アニメにおける声優の音声は「話者の声でもあり、キャラクターの声でもある」声だと考えることができ、話者とキャラクターが重ね合わせ状態で存在しているととらえることが可能である。そしてだからこそ、「キャラクターと声優を同一視する——キャラクターを観るときに声優のイメージをそのキャラクターに重ね、逆に声優を観るときにキャラクターのイメージをその声優に重ねる——という認識の基盤が、少なくともアニメ声優に関心を寄せるファンの間ですでに共有されている(程, 2023)」という事態が生じると考えられる。

また、日本のアニメでは女性声優が少年キャラクターを演じることが多く(石田, 2020)、そのような状況では、話者とキャラクターの重ね合わせは、男性と女性との重ね合わせでもある。すなわち、女性声優の少年声は「女性でもあり、男性でもある」声としてとらえることができると考えられる。

3.2.2 二・五次元文化における声

前項で言及したアニメにおける声優の声は、基本的に声優自身の身体は見えない状態での声であり、「声優の声でもありキャラクターの声でもある」ような重ね合わせ状態が、比較的起こりやすい状況だ

と考えられる。では、声優の身体が聞き手に見えている状態で、声優がキャラクターとして発した音声は、アニメの場合と同様に重ね合わせ状態を生じうるのだろうか。

声優について、「その技能が専門化していく中でキャラクター文化と結びつき、メディアミックスハブとしての役割を声優が担っていくことになる(永田, 2024)」と指摘されており、現代の声優の仕事は「キャラクターを様々なメディアで演じること」となっている。様々なメディアとは、ラジオやイベント登壇、メディア取材、ゲーム、派生作品なども存在するが、その1つとして、声優が「キャラクターとして」舞台上に立ち、キャラクター名義の曲を歌ってライブを行っているような、声優/キャラ・ライブコンサートが存在している。そのようなコンサートは、「マンガ、アニメ、ゲームなどの視覚情報が優先される二次元の虚構世界と、身体性を伴う経験を共有する三次元の“現実”世界のあいまいな境界でおこなわれる文化実践を、『二・五次元文化』として定義する(須川, 2021, p. 16)」ならば、二・五次元文化の1つとしてとらえることができる。ここでの二・五次元は「二次元の虚構キャラクターが、三次元の現実で具現化し、虚構性を限りなく内包した人間の生身の身体、いわば『虚構的身体(“virtual corporality”)』と呼びうるものによって実践されること(須川, 2021, p. 15)」であり、声優/キャラ・ライブコンサートはまさにその好例だと考えられる。そして声優/キャラ・ライブコンサートにおいて、「声優は『キャラ』の模倣ではなく、『キャラ』そのものの生成として、オーディエンスによって経験される(川村, 2024, p. 107)」。そしてその際に、「オーディエンスの情動を最も強く触発する記号性は、いうまでもなく声」であり、『キャラ』の分配的領域を構成する声が声優によって発せられることで、『キャラ』と声優を不可分のものとみなすリアリティが横断個体的に経験される(川村, 2024, p. 107)」ことになる。これは「声を媒介にした声優の身体とキャラクターの一体化(須川, 2021, p. 118)」であり、「登場人物と同じ声色による歌声が聞こえてきた際には、声優が歌っているのではなく、登場人物が歌っているという想像が生じることがある(シノハラ, 2017, p. 86)」。そしてそのような事態の一例として、藤田茜が演じた荒木レナというキャラクターのライブについて、「舞台上で発せられた『跳ぶよ』を、荒木レナの言葉として(虚構的に)聞くか、藤田茜の言葉として聞くかは

1 未確定でもあり、観客に委ねられている（シノハラ、
2 2017,p.93）」と述べられている。すなわち、声優自身
3 の身体が見えない状態に限らず、声優の身体が聞き
4 手に見えている状態であったとしても、声優がキャ
5 ラクターとして発した音声は「話者の声でもあり、キ
6 ャラクターの声でもある」声であり、話者とキャラク
7 ターが重ね合わせ状態で存在していると考えること
8 ができる。なお、視聴者は役者とキャラクターをごち
9 や混ぜにする傾向があることも示されており
10 （Goldstein & Filipe, 2018）、このような重ね合わせ状
11 態は声優に限らず、二・五次元一般、あるいは「キャ
12 ラクターを演じながら声を発している人」を見る状
13 況一般で生じうるものであるかもしれない。

14 3.2.3 VTuber の声

15 本項では、前項までの内容を踏まえつつ、昨今隆盛
16 を見せるバーチャル YouTuber (VTuber) の声につい
17 て検討を行う。VTuber も「二・五次元文化の一例（須
18 川, 2021, p. 15）」として捉えることが可能である。し
19 かし、前項で述べた声優/キャラ・ライブコンサート
20 のような状況とは異なり、VTuber は動画サイト上で
21 の配信であってもライブであっても、基本的に「中の
22 人（配信者・演者）」の物理的身体が聞き手の目にさ
23 らされることはなく、「パーソンの身体は後ろに隠れ、
24 ペルソナの身体だけが鑑賞者の前に現れている（難
25 波, 2018, p. 119）」（もちろん、例外はあるものの）。そ
26 のため、前項の二・五次元文化における声とは項を分
27 けて考えていく。ここでの問いは、「鑑賞される
28 VTuber とは誰だろうか。たとえば、VTuber が笑った
29 というとき、誰が笑っているのだろうか。VTuber が
30 愛らしいというとき、誰が愛らしいのだろうか（難波,
31 2018, p.117).」というもの、すなわち VTuber が声を
32 出すとき、誰が声を出しているのか、あるいは誰の声
33 として認知されているのであろうか、というもので
34 ある。

35 VTuber のあり方は非常に多様であり、一意に定め
36 て検討することは難しい。この点は、「いっぽうで、
37 ある特定の VTuber は、その外見から一般に想像しう
38 るような声や性格を逸脱し、かつ、そのパーソンの経
39 験や考えをあからさまに語っているようにみえる。
40 たほうで、動画内にとどまらず、SNS においておロ
41 ールプレイに徹しており、あくまでもキャラクタと
42 して鑑賞者に鑑賞される VTuber もおり（難波, 2018,
43 p. 120）」などと指摘されている。ただし、ここまでに
44 主として述べたアニメとの比較を考えると、アニメ
45 における声優の演技はあくまで「作品世界における

46 キャラクターを演じる」ものであるのに対して、多く
47 の場合で VTuber は、作品世界の中だけで生きている
48 のではなく、視聴者の生きるこの物理世界と何らか
49 の意味で地続きの世界で生きているように思われる。
50 それはたとえば「月ノ美兎は水を飲む」と題した論考
51 が存在することや（新, 2018）、VTuber の存在論に関
52 する論考の冒頭に「月ノ美兎は、カフェオレを飲んで
53 いた（篠崎, 2024, p. 271）」「紫咲シオンが、とある深
54 夜配信中にお腹が空いてカレーメシを食べ始めた出
55 来事を考えてみよう（本間, 2024, p. 323）」のような文
56 が存在することからも示唆されている。この点を踏
57 まえて、「VTuber は生身の人間である中の人（配信者・
58 演者）と同一の存在である」と考えることも可能であ
59 る（篠崎, 2024）。その場合、VTuber の声については
60 「もはやこれらのキャラクターが喋るとき、その向
61 こうに本当の発声者がいることを、誰もが疑ってい
62 ない」「彼女らの声は…あくまでもそれぞれ「個
63 人」の身体に固定された声であるかのように振る舞
64 っている（黒寄, 2018, p. 195）」ととらえることができ
65 だろう。「実質的には演者は顔を出していないのに
66 も拘わらず、キャラクターの顔によって「声」を発し
67 ている存在が視聴者に明確に見える状態だと思って
68 もらえる（皇牙, 2018, p. 65）」という記述も合わせて
69 考えると、「VTuber の声は、中の人（配信者・演者）
70 の声である」と考えることも可能だと考えられる。こ
71 のような状況は、ある意味で VTuber というキャラク
72 ターの声と、中の人（配信者・演者）という話者の声
73 が重ね合わせ状態にあるととらえることもできるか
74 もしい。

75 他方で、「VTuber は生身の人間である中の人（配信
76 者・演者）と同一の存在である」とは必ずしも言い切
77 れない。視聴者は中の人（配信者・演者）と VTuber
78 を分けて見ており（Lu et al., 2021）、少なくとも視聴
79 者、すなわち声の聞き手にとっては「演者とキャラク
80 ターは別物だ」という考えが主流（皇牙, 2018, p. 65）」
81 だとも考えられている。「パーソンそのものは、しか
82 し、直接わたしたちとコミュニケーションをとって
83 いるわけではない（難波, 2018, p.118）」のである。こ
84 の点について、アニメにおける声優とも関連付けて
85 検討できる考え方として、「技術的所産」として
86 VTuber をとらえた場合に、「現在の典型的な VTuber
87 の場合、キャラクターの想像に対して配信者のなす
88 貢献が極めて大きい…その話し声や歌声は基本的に
89 配信者の貢献分である（松本, 2024, p.318）」と考える

1 ことができる。この場合の VTuber の声は、3.2.1 項で
2 述べたアニメにおける声優の声と同じような形で、
3 話者とキャラクターの重ね合わせ状態として理解で
4 きるかもしれない。ただし前述したとおり、アニメに
5 おける声優の声と VTuber の声は「作品世界における
6 キャラクターの声か否か」という違いが一般には存
7 在するため、3.2.1 項の議論をそのまま当てはめるこ
8 とができるかについては留意が必要である。

9 いずれにしても、『声』は、音声コミュニケーションが
10 メインになるメタバースの世界においても
11 VTuber として配信するにあたって、アイデンティ
12 ティの印象を左右する決定的な要素（関根, 2024,
13 p.186）」であり、「私たちは、配信者の声だけが吹き
14 込まれている実写動画に対しても、それを VTuber の
15 声として鑑賞することが可能である（山野, 2024, p.
16 83）」ことを考えても、VTuber という存在において声
17 は 1 つの重要な要素である。そのため、VTuber の声
18 における重ね合わせを検討することは、声における
19 重ね合わせ状態の理解にとって有用な 1 つの観点で
20 あると考えられる。

21 4. 声における「人間らしさ」

22 我々は声を聞いて、その背後に人間を想像する。
23 「人間は声の向こうに誰かの身体を、そして名前を
24 もつ一つの人格を想定せずにいられない（増野, 2014,
25 p.217）」と語られる通り、発する「人間」と声は深く
26 結びついたものであるが、音声合成技術の発展に伴
27 い、人間ではなく機械が発する音声が日常的に使わ
28 れるようになってきている。本章では、声と人間らし
29 さとの関係についての研究を概観した上で、声にお
30 いて人間と機械がいかに重なりうるのかについて検
31 討を行う。

32 4.1 「人間らしさ」の認知

33 声の人間らしさを定義するのは難しいものの、1 つ
34 の指標として「読み上げの自然さ」を挙げることがで
35 きる（三井, 2010）。これは聞き手にとっての「言語内
36 容の聞きやすさ」と関わると思われ、音声合成による
37 機械の声を人間の声に近づける、というアプローチ
38 を考えた場合、まずは読み上げの自然さについて、人
39 間と遜色ないものを目指すことが考えられる。以前
40 は人間の肉声に比べて機械音声はこの点が劣ってお
41 り、たとえば館内放送を想定した音声を用いた研究
42 において、肉声の方が聞きやすいことが示されてい
43 た（三谷, 2014）。しかし、技術の進歩は著しく、十分

44 なデータ量と計算量があれば、人間と遜色のない、す
45 なわち主観評価が肉声と統計的に変わらない機械音声
46 を合成することは、すでに技術的には可能となって
47 いる（Tan et al., 2024）。

48 では、すでに十分に「人間らしい」機械音声が出来
49 上がっているかといえば、依然として課題は残って
50 いる。たとえば、読み上げの自然さだけでなく、感情
51 をうまく機械音声に乗せることも、人間らしい機械
52 音声を実現する上では重要な課題だと考えられる
53 （Qi et al., 2024）。また、対話場面を考えると、実際
54 の人は親密さによって声や話し方が変わっていくた
55 め（Chiba & Ito, 2024）、そのような変化を反映させる
56 ことも、自然な対話音声を実現するには重要かもし
57 れない。関連して、共感的な対話音声の生成や（Saito
58 et al., 2022）、笑い声の合成（Xin et al., 2023）など、
59 より自然な人間の発話に近い機械音声の合成に向け
60 た研究が進められている。朗読的な発話よりも自発
61 的な発話をする機械音声で話すエージェントの方が、
62 人間との会話に近いと評価されることも示されてお
63 り（Iizuka & Mori, 2022）、そのような自発音声合成に
64 ついても研究が行われている（Matsunaga et al., 2023）。
65 さらに、人間相手に話す状況だけでなく、AI 同士で
66 話している状況で、より自然な人間のように感じる
67 対話を目指した研究も行われている（Mitsui et al.,
68 2023; Nguyen et al., 2023）。話者が人間の場合と機械
69 の場合で比較した研究をみると、人間の声と機械の
70 声いずれに対しても、感情的な発話の模倣が起こる
71 ことを示した研究や（Cohn et al., 2024）、人間の歌声
72 と機械音声の歌声を比較した際に、定型発達者だと
73 人間の歌声をよりポジティブに評価する一方で、自
74 閉スペクトラム症の方では人間の歌声と機械音声の
75 歌声で評価に差がないことを示した研究などが存在
76 している（Kuriki et al., 2016）。このようにまだ課題は
77 存在しているものの、「人間らしい」機械音声につい
78 ては様々な研究が進展しており、将来的には人間と
79 区別のつかない対話が可能な機械音声的合成できる
80 ようになると考えられる。

81 全く区別のつかない状態となる前段階では、肉声
82 と機械音声の使い分けも重要となる。たとえば、相手
83 が人間の場合と機械の場合で、話しかけ方や話す内
84 容が違うことが示唆されており（山田他, 2005）、適切
85 な使い分けが必要となるだろう。また、同じ機械音声
86 であっても、自動電話応答場面とキャラクタ音声合
87 成サービスではユーザに求められるものが変わって

1 くることも報告されており，前者では明瞭性や韻律
 2 の自然性，肉声感などが重要となる一方で，後者では
 3 それらのある種の「人間らしさ」はあまり重要ではな
 4 く，キャラクターや韻律制御の自由度が重要となる
 5 ことが示されている（三井, 2010）．このように，人間
 6 の声と機械の声を場面に合わせて適切に使い分けな
 7 がら，徐々に機械の声は社会実装されていき，そして
 8 将来的には，聴覚印象として人間の声と機械の声は
 9 区別がつかない状態になっていくと考えられる．

10 4.2 人間と機械の重ね合わせ

11 前節において概観した，声における「人間らしさ」
 12 の認知に関する研究を踏まえつつ，本節では，声にお
 13 いて「人間」と「機械」の重ね合わせが生じる可能性
 14 について検討を行う．

15 すでに社会の様々な画面で機械音声は活用されて
 16 おり，駅などでのアナウンス音声で機械音を聞く
 17 機会も増えてきている．また，動画サイト等では合成
 18 音声キャラクターを用いた動画も多く投稿されてい
 19 る．現状，これらの音声は多くの場合，人間と区別が
 20 つかない状態までは達しておらず，その意味で「人間
 21 の声でもあり，機械の声でもある」という解釈が重な
 22 った状態は生じていないと考えられる．ただしここ
 23 で，重ね合わせ状態とは直接関わらないものの，機械
 24 音声においては人間の声とは異なる形で「キャラク
 25 ターと声」の関係が生じていることを指摘したい．機
 26 械音声とキャラクターとの関係について，最も多く
 27 語られているのは歌声合成ソフトであるボーカロイ
 28 ド，その中でも初音ミクである．初音ミクの声は，声
 29 優の藤田咲の声をサンプリングしたことが明示的に
 30 知られている．しかし，初音ミクの声は，藤田咲の声
 31 として知覚されるわけではなく，「ボーカロイドの素
 32 材は声優や歌手が提供した人間の声であるが，元の
 33 持ち主の身体からは，はっきりと切り離されている
 34 （増野, 2014, pp. 216-217）」．そのように人間の声とし
 35 ては知覚されない初音ミクの声は，他方で絵を含め
 36 たキャラクターとは重なる存在であり，アニメを含
 37 め多くの場合に先に絵があった上でそこに声が当て
 38 られるのに対して，「初音ミクが声と絵という二つの
 39 要素から成り立つもので，とりわけ前者の存在を前
 40 提としながら後者が要請されていた（さやわか, 2015,
 41 p.26）」ことは，キャラクターと声の関係を考える上
 42 で重要である．また，初音ミクの声は確かに藤田咲の
 43 声としては認知されないものの，何らかの繋がりを持
 44 っていることは考えられ，「藤田咲さんの声が断片

45 化されて，ちょっと違う形になって初音ミクとして
 46 出てくるわけなんですね．でも，元の声にあった『私，
 47 可愛いでしょ』という雰囲気は，ある種ランダム化さ
 48 れて残っている（柴, 2014, p. 105）」のような指摘が存
 49 在している．さらに，初音ミクは「人間ではない」か
 50 らこそ，人間が声を発する場合以上に，性別を超えら
 51 れる可能性があり，「ミクは女性で女声だが，音源を
 52 低音に操作すれば男声にすることもできる…男であ
 53 り，女．女であり，男（村瀬・須川, 2014, p.91）」と
 54 の言説が存在している．

55 いずれにせよ，ここまでの検討では，「人間の声で
 56 もあり，機械の声でもある」という解釈が重なった状
 57 態は生じていない．しかし前節で述べたように，研究
 58 レベルでは人間と遜色のない機械音を合成するこ
 59 とは技術的に可能となっており（Tan et al., 2024），将
 60 来的にはそのような機械音声为社会実装されると考
 61 えられる．その際，物理的な身体も伴わせようと思
 62 うと，人間と区別のつかない物理身体も生み出す必要
 63 があり，技術的なハードルは高くなる．一方で，声の
 64 みであったり，バーチャル環境であったりすれば，こ
 65 の点は容易にクリアできる．たとえば CoeFont の「お
 66 しゃべりひろゆきメーカー」や Cotomo の「緑川光の
 67 声」のように，タレントや声優などの有名人の声をサ
 68 ンプリングし，別のキャラクターではなく「その人の
 69 名前」を用いた機械音声が存在しているが，それらの
 70 音声は人間と遜色がつかない状態となった場合，事
 71 前に知らされなければ，その有名人の肉声なのか機
 72 械音声なのか区別がつかない状態が生じると考えら
 73 れる．そのような声は，少なくとも聞き手にとっては
 74 「機械の声」でありながら「人間の声」でもある状態，
 75 すなわち人間と機械の重ね合わせ状態となる．また，
 76 たとえばソーシャル VR プラットフォーム上で，自
 77 律的に会話できる AI エージェントが実装された際に
 78 （倉井・平木, 2023），その音声は人間と区別がつか
 79 なければ，話し手の情報を知らない聞き手にとって，コ
 80 ミュニケーション相手は「機械だが，人間でもある」
 81 状態となるだろう．さらに新しい形として，人間であ
 82 る中の人（配信者・演者）が存在する VTuber とは異
 83 なり，背後に人間がいない AITuber が動画サイト等
 84 で配信することが行われている．AITuber はモデルの
 85 見た目だけを取り出すならば，VTuber と大きな違い
 86 はないため，その存在が「人間らしく」なる上では，
 87 モデルの動きに加えて，声の人間らしさが重要な要
 88 素になると考えられる．AITuber のある配信について，

1 「それまでの『AITuber』の発話は、まだ『機械』の
2 感じが強く残っていたのですが、今回アップデート
3 されたアイリーンさんの発話は、自然な人間の発話
4 にかなり近いものになっていた」ことに基づき
5 「AITuber は実在の配信者と実質的に同じように存
6 在しうる（すなわちバーチャルに存在しうる）という
7 可能性を感じました（山野, 2025, p. 78）」と述べた言
8 説も存在しており、技術の進歩に伴って、聞き手にと
9 って「機械の声であるが、人間の声でもある」という
10 解釈の重ね合わせ状態は、より一般的に生じる事例
11 となっていくのかもしれない。

12 5 今後の展望

13 本稿ではここまで、声における「私らしさ」「その
14 人らしさ」「人間らしさ」という3つの観点から先行
15 研究を概観した上で、それぞれの観点で「自己と他者」
16 「話者とキャラクター」「人間と機械」の重ね合わせ
17 状態が生じている可能性について検討してきた。そ
18 れを受けて本章では、声における「私らしさ」「その
19 人らしさ」「人間らしさ」それぞれと紐づく問いを提
20 起しつつ、主に認知科学の観点からどのような研究
21 の方向性が考えられるかについて議論を行う。

22 5.1 声における「私らしさ」に紐づく問い

23 声における「私らしさ」と重ね合わせについては、
24 日常場面も含めた「声で演じる」行為との関連もあり
25 つつ、声質変換技術の進展による「自己」と「他者」
26 の重ね合わせ状態が生じている可能性が示された。
27 これは、声質変換技術により、新たな自己を形成して
28 いける可能性を示唆している。音声の基本周波数を
29 操作してフィードバックすることで感情状態を誘導
30 できることを示した研究や(Aucouturier et al., 2016),
31 自信のある状態の声をフィードバックすることによ
32 る緊張緩和の可能性を示した研究もあり(成瀬他,
33 2019), 自身の発した声を操作することで、話者の状
34 態はコントロールできる可能性がある。それは「自己」
35 という感覚でも同様であり、キャラクター性の獲得を
36 はじめ(倉田他, 2021), 声をデザインすることで新た
37 な自己をデザインすることが可能となっていくかも
38 しない(Arakawa et al., 2021; 高道, 2021)。その際、
39 理想の自己を表現する声を用いることも可能だろう
40 し、逆に画一的な声を使ってアイデンティティを「隠
41 す」こともできる(Hatada et al., 2023)。「『声』は…コ
42 ミュニケーションを通して私たち自身の存在を示す、
43 音響世界における『アバター』とも言えるべき、重要な

44 アイデンティティの一つ（バーチャル美少女ねむ、
45 2022, p.178)」であり、技術の発展により、「自己」と
46 という感覚はある意味でゆらぎ、そして新たに形成さ
47 れていくと考えられる。そのような状況において、声
48 が変わることゆえに「私」は、いかに織り上げられ
49 てナラティブとしての私になっていくのであろうか、
50 という問いは、これからの時代において重要になっ
51 てくるであろう。

52 そしてそのように、望む自己をデザインしていけ
53 るかもしれない可能性がある一方で、「声はヒトの身
54 体と深く一体化した圧倒的な存在感をもつ音（増野,
55 2014, p.12)」でもある。すなわち、私たちは物理的な
56 身体を有しており、それは声を変えたとしても変わ
57 るわけではない。声を変えても変わらない、あるいは
58 変えられない、身体と紐づいたままならぬ「私」はい
59 るのか、という問いについても、合わせて検討してい
60 くことが、声における「私らしさ」を包括的に考えて
61 いく上で、ひいては「自己とは何か」を改めて考えて
62 いく上で重要だと考えられる。

63 5.2 声における「その人らしさ」に紐づく問い

64 声における「その人らしさ」と重ね合わせについて
65 は、アニメや二・五次元文化、VTuberなどのポップ
66 カルチャーにおける人間の声の技術によって、「話者」
67 と「キャラクター」の重ね合わせ状態が聞き手の中で
68 生じている可能性が示された。このような重ね合わ
69 せ状態は、声に対する「好き」という感覚についての
70 問いを喚起している。声のよさや魅力については
71 様々な研究が行われており(Weiss et al., 2021), 声の
72 高さとの関係や(Jones et al., 2010), 平均すると魅力
73 が高まるという研究(Bruckert et al., 2010), 他方で平
74 均が必ずしも魅力を高めるわけではないという研究
75 (Ostrega et al., 2024)などが存在している。よさや魅
76 力に関する評価の個人差についても、評価の性差と
77 音響特徴量との関係についての研究(横森他, 2016)
78 や、性差や年代差も含めて大規模な音声・参加者を対
79 象に行った研究が存在している(Suda et al., 2024)。
80 これらの研究では基本的に、ある声は「特定の話者の
81 声」であり、そこに重ね合わせは想定されていない。
82 しかし3.2節で検討したように、特にポップカルチャー
83 においては声において話者とキャラクターの重ね
84 合わせが生じているとすると、声を聞いて心惹かれ
85 て好きだと感じた場合に、私たちは「誰の」声に心惹
86 かれているのだろうか？ポップカルチャーにおける
87 声について、「繰り返し聴取は、声の演戯性が持つ声

1 そのものの魅力へと没頭させ、まさに声への『偏愛』
 2 『中毒現象』を生じさせる…声優によって発せられ
 3 る声は、一方で、その実在性により架空のキャラクタ
 4 ーの想像をヴィヴィッドなものとし、他方で、その演
 5 戯性によりキャラクターへの愛着を持たせている
 6 (シノハラ, 2017, p.78)」と指摘されているが、その
 7 「偏愛」や「愛着」のようなこだわりの対象は、声優
 8 の声なのであろうか、あるいはキャラクターの声な
 9 ののであろうか。

10 また、3.2 節では女性声優の少年声について「女性
 11 でもあり、男性でもある」声として言及している。「声
 12 優の受容には二つの傾向が存在する。一つは、年齢と
 13 性別、さらには外見上の声優とキャラクターの乖離
 14 を認識しながら、声優の演技を楽しむ傾向である。も
 15 う一つは、年齢と性別、さらには外見上のキャラクタ
 16 ーと声優の一致を求めながら、声優の演技を楽しむ
 17 傾向である(石田, 2020, p. 197)」との指摘を踏まえた
 18 際に、女性声優の少年声は「前者の中核に位置する
 19 (石田, 2020, p. 197)」と考えられる。男性と女性は一
 20 見相反するものであるため、このような状態が生じ
 21 ることは、単純な話者とキャラクターの重ね合わせ
 22 状態以上に、矛盾や曖昧性を声が内包しうることを
 23 示している。そしてそのように解釈が定まらず、複数
 24 の解釈が並列している不定性があるからこそ生じる
 25 美的体験の存在が示唆されており(布山・西郷, 2022),
 26 声における「その人らしさ」と声に対する「好き」と
 27 という感覚との関係を十分に理解する上では、声が内
 28 包しうる矛盾や曖昧性、不定性があるからこそ生じ
 29 るよさや魅力とはどのようなものか、という問いに
 30 ついても検討していく必要があるだろう。

31 5.3 声における「人間らしさ」に紐づく問い

32 声における「人間らしさ」と重ね合わせについては、
 33 音声合成技術の発展に伴って「人間」と「機械」の重
 34 ね合わせ状態が生じることが示された。すなわち、
 35 コミュニケーションの相手が人間であるか機械であ
 36 るか、少なくとも声だけでは区別がつかない状態が
 37 訪れると考えられるが、そのような状況において、私
 38 たちは「誰と」コミュニケーションを通じてつながり
 39 たいと感じるのであろうか。たとえば絵画を用いた
 40 研究では、実際にはどちらも機械が描いた絵画であ
 41 っても、「人間が描いた」とラベルづけされた絵画の
 42 方をポジティブに評価することや(Bellaiche et al.,
 43 2023), 人間の描いた絵画と機械の描いた絵画を区別
 44 することはできなくとも、なぜか人間の描いた絵画

45 の方がポジティブに評価されたことが報告されてい
 46 る(Samo & Highhouse, 2023)。このように「人間の方
 47 がいい」と評価されるケースもある一方で、たとえば
 48 ロボットを用いた研究では、日本人の若年女性だと
 49 ロボットでも人間でも自己開示の程度が変わらない
 50 ことが示されていたり(Uchida et al., 2020)、高齢者
 51 のライフレビューでは話し相手がロボットの方がポ
 52 ジティブとなりうる可能性が示されていたりする
 53 (Ueda et al., 2021)。また、ボーカロイドの歌声は機
 54 械の歌声であるが、「ボカロを抵抗なく聴くリスナー
 55 も、それを声と等価なものとしてではなく、全く別物
 56 として受容している(宮本, 2017, p.36)」との指摘が
 57 あり、人間と比較するのではなく、機械の声は機械の
 58 声として楽しむことも考えられる。すなわち「いつで
 59 も人間がよい」わけではなく、逆に「いつでも機械が
 60 よい」わけでもないと考えられるため、つながる相手
 61 としての人間と機械は何が似ていて何が異なってい
 62 るのか、という問いがこれから重要になってくると
 63 考えられる。

64 そして私たちは「声の向こうに誰かの身体を、そし
 65 て名前をもつ一つの人格を想定せずにいられない
 66 (増野, 2014, p.217)」としたら、人間と機械の区別が
 67 つかなくなった世界において、私たちは声の“背後”
 68 にどのような存在を感じ、どのようにつながりたい
 69 と思うのであろうか。相手がいて初めて成立するコ
 70 ミュニケーションという場において、話し相手が人
 71 間か機械かも分からない状態で、私たちは声の“背後”
 72 にいると感じる存在とつながっている、あるいは理
 73 解し合えていると感じることができるのであろうか。
 74 このような声における「人間らしさ」に紐づく問いは、
 75 改めて「人間とは何か」について考える契機となるだ
 76 ろう。

77 5.4 認知科学研究のアプローチ

78 前節まで、声における「私らしさ」「その人らしさ」
 79 「人間らしさ」それぞれと紐づいて喚起されうる問
 80 いを検討してきた。それらを踏まえつつ本節では、主
 81 に認知科学の観点から、これらの問いに対してどの
 82 ような研究の方向性が考えられるかについて検討を
 83 行う。

84 まず本稿では、声における様々な重ね合わせ状態
 85 の可能性を考えてきたが、いずれの場合においても、
 86 対象となる「声」は物理的には1つである。つまり生
 87 じている重ね合わせは、ヒトの認知における解釈の
 88 重ね合わせであり、認知科学の研究対象となる。量子

1 論的世界観を踏まえた認知科学研究として、たとえ
2 ば文章理解における複数解釈の重ね合わせ状態につ
3 いて、量子確率論の公理に基づいて数理モデルで表
4 現した研究があり(布山・西郷, 2022), そのようなア
5 プローチを用いることで、主として人文社会学的な
6 様々な論考において言及されてきた声における重ね
7 合わせ状態について、異なる観点からの研究が可能
8 となるかもしれない。

9 また、本稿で検討した声における重ね合わせ状態
10 は、認知科学の概念であるプロジェクション・サイエ
11 ンスと相性がよいように思われる。たとえば自己と
12 他者の重ね合わせについては、アバターによる自己
13 変容と関連付けて議論されているゴーストエンジニア
14 アリングと紐づけて検討ができるであろう(鳴海,
15 2019)。話者とキャラクターの重ね合わせについては、
16 たとえば二・五次元舞台についてプロジェクション・
17 サイエンスの観点から検討がなされており(久保,
18 2022)、プロジェクション・サイエンスの領域で蓄積
19 されてきた知見を活かした研究が考えられる。また、
20 主に話者とキャラクターとの重ね合わせ状態に関す
21 る検討と関連する、哲学・美学領域でフィクション鑑
22 賞の理論として用いられるメイクビリーブ(想像力)
23 に関する議論は、プロジェクションという概念との
24 親和性が高いと考えられる(松本, 2017; ウォルトン,
25 2016)。想像の引き金として働く事物がプロップ(小
26 道具)であり、アニメにおける声優の声は、キャラク
27 ターの「声についての虚構的真理を生成するプロッ
28 プ(松本, 2017, p.106)」である。そして、「虚構上の
29 キャラクターである荒木レナに対して、実在する人
30 間である藤田茜の声があてられることで、少なくと
31 も音声面について、荒木レナの想像には実体が与え
32 られる…藤田茜について、その声が荒木レナの声で
33 あると想像(シノハラ, 2017, p.70)」するとあるよう
34 に、話者とキャラクターの重ね合わせ状態が生じる
35 上で、この想像力に関する議論は重要である。そして
36 このようなフィクション鑑賞における想像力に関す
37 るの理論は、プロジェクション・サイエンスにおける
38 「表象を特定の人や事物に投射(プロジェクション)
39 する」こととの相性がよいように思われ、これらの議
40 論を接続することで、声における重ね合わせ状態に
41 ついてより深い理解が得られるかもしれない。

42 先述したように、声における重ね合わせ状態はヒ
43 トの認知・解釈によって生じるものではあるが、重ね
44 合わせ状態となりやすい「声」は存在している可能性

45 がある。より具体的には、ポップカルチャーにおける
46 人間の声の技術を考えた際に、たとえば「女性声優の
47 少年声」には、通常の少年の声や他の演技音声とは異
48 なる特有の「女性声優の少年声らしさ」に関する音響
49 特徴が内包されている可能性がある(林・森勢, 2025)。
50 また、アニメにおける声優の声を考えても、「アニメ
51 らしい声」についての調査を行った研究や(林・大杉,
52 2021)、実際のアニメにおける声質を調べた研究が
53 様々な存在している(原・伊藤, 2009; Ishi et al., 2023;
54 Starr, 2015; Utsugi et al., 2019)。あるいは、自発的な発
55 話と演技発話の間で音響特徴量が異なることについ
56 ても報告されており(Jurgens et al., 2011; Laukka et al.,
57 2011; Scherer, 2003; Williams & Stevens, 1972)、アニメ
58 やキャラクター、演技はいずれも「声」に通常の自発
59 発話とは異なる特徴を持っていると考えられる。そ
60 うだとすると、たとえば「重ね合わせ状態にある声を
61 情報科学的に認識する研究」や、「重ね合わせ状態に
62 ある声とそうでない声の聴取印象の違いに関する研
63 究」なども、方向性としては考えられるかもしれない。

64 5.5 技術と社会との関係

65 本稿で検討した声における重ね合わせ状態は、主
66 にポップカルチャーにおける人間の声の技術と、音
67 声合成や声質変換などの音声工学技術と深く関わる
68 ものであった。今後の展望を考える上では、それぞれ
69 の技術と社会との関係についても検討することが重
70 要である。

71 まず人間の声の技術については、「今あるものを、
72 未来にどのように残していくのか」というアーカイ
73 ブの観点が重要となってくるだろう。研究用の音声
74 データベースやコーパスは多数存在しており、また
75 歴史音声のデジタルアーカイブに関する検討も行わ
76 れているが(Saeki et al., 2023)、ポップカルチャーに
77 おける人間の声に関するデジタルアーカイブについ
78 ては、メディア芸術に関するデジタルアーカイブの
79 取り組みの中でも十分に言及されていないのが現状
80 である(三原, 2022)。声における重ね合わせ状態につ
81 いての研究を行う上でも、また文化を守り伝えてい
82 く上でも、人間の声の技術のデジタルアーカイブに
83 ついて検討していくことは重要である。

84 また、音声工学技術については、新たな科学技術の
85 社会実装を考える際に必ず検討すべき点として、倫
86 理的・社会的・法的課題(ELSI)についての議論が不
87 可欠である(標葉・見上, 2024)。音声工学技術のELSI
88 については、人の容姿における肖像権のように、声の

1 人格権に関する法的な検討が行われていたり（荒岡
2 他, 2023）、論文のシステマティックレビューを通じ
3 てELSIに関する論点の抽出を試みた研究が存在した
4 りするものの（Barnett, 2023）、十分な議論がなされて
5 いるとは言い難い状況であり、産官学民を巻き込ん
6 ださらなる検討が重要となるだろう。

7 5.6 本論文の限界

8 本稿の限界として、声と深く関係するにもかかわ
9 らず、「言語」「音楽（歌声）」「多感覚（見た目など
10 の関係）」に十分に言及できなかった点を挙げたい。
11 いずれも声について研究する上で非常に重要な要素
12 であるにもかかわらず、一部言及されている場合は
13 ありつつ、本稿では十分な検討が行えていない。それ
14 はこれらの要素が重要でないことを意味するもので
15 は全くなく、本稿が「声単体の」「歌声ではない発話
16 音声における」「広義の声質（音声波から知覚される
17 韻質以外の聴覚上の特質（森他, 2014; 日本音響学会,
18 2003, p. 194)）」について主に検討したこと起因す
19 る。声における重ね合わせ状態についての研究を進
20 めていく中で、これらの要素については別の機会に
21 整理を行っていきたい。

22 6 おわりに

23 本稿では声における「私らしさ」「その人らしさ」
24 「人間らしさ」それぞれと関係する重ね合わせ状態
25 について検討し、今後認知科学を始めとした諸領域
26 において、どのような研究が行えるかについて展望
27 を述べた。検討を通じて提起された「私らしさ」「そ
28 の人らしさ」「人間らしさ」と紐づく問いは、いずれ
29 もこれからの社会における新たな「自己観」や「人間
30 観」と繋がりをうめるものである。このことは、量子論的
31 世界観に基づいて、声における「自己と他者」「話者
32 とキャラクター」「人間と機械」の重ね合わせ状態に
33 ついて検討することで、未来の自己観や人間観がど
34 のように形作られていくのかについて検討できるこ
35 とを示唆している。また本稿で述べた議論は、人間の
36 技術（アート）と工学技術（エンジニアリング）、あ
37 るいは人文社会系と理工系など、対比されやすいも
38 のを混ぜ合わせて検討することの価値を示唆してい
39 る。声を題材とした研究の中で、様々な“間”を知り、
40 繋いでいくことで、未来の日常における自己や人間
41 についての理解が深まっていくことに本稿が寄与で
42 きればと願っている。

43 文 献

- 44 Arakawa, R., Kashino, Z., Takamichi, S., Verhulst, A., & Inami,
45 M. (2021). Digital speech makeup: Voice conversion based
46 altered auditory feedback for transforming self-
47 representation. *Proceedings of the 2021 International
48 Conference on Multimodal Interaction*, 159-167.
- 49 荒岡 草馬・篠田 詩織・藤村 明子・成原慧 (2023). 声の人
50 格権に関する検討 情報ネットワーク・ローレビュー,
51 22, 24-44.
- 52 Aronovitch, C. D. (1976). The voice of personality: stereotyped
53 judgments and their relation to voice quality and sex of
54 speaker. *Journal of Social Psychology*, 99(2), 207-220.
- 55 浅井 智久・丹野 義彦 (2010). 声の中の自己と他者：幻聴
56 の自己モニタリング仮説 心理学研究, 81(3), 247-261.
- 57 新 八角 (2018). 月ノ美兎は水を飲む ユリイカ, 50(9), 92-
58 99.
- 59 Aucouturier, J. J., Johansson, P., Hall, L., Segnini, R., Mercadié,
60 L., & Watanabe, K. (2016). Covert digital manipulation of
61 vocal emotion alter speakers' emotional states in a congruent
62 direction. *Proceedings of the National Academy of Sciences*,
63 113(4), 948-953.
- 64 バーチャル美少女ねむ (2022). メタバース進化論 仮想現
65 実の荒野に芽吹く「解放」と「創造」の新世界 技術評
66 論社
- 67 Barnett, J. (2023). The ethical implications of generative audio
68 models: A systematic literature review. *Proceedings of the
69 2023 AAAI/ACM Conference on AI, Ethics, and Society*,
70 146-161.
- 71 Bellaiche, L., Shahi, R., Turpin, M. H., Ragnhildstveit, A.,
72 Sprockett, S., Barr, N., Christensen, A., & Seli, P. (2023).
73 Humans versus AI: whether and why we prefer human-
74 created compared to AI-created artwork. *Cognitive
75 Research: Principles and Implications*, 8:42, 1-22.
- 76 Belin, P., Boehme, B., & McAleer, P. (2017). The sound of
77 trustworthiness: Acoustic-based modulation of perceived
78 voice personality. *PLoS ONE*, 12(10), e0185651.
- 79 Bredikhina, L. (2022). Babiniku: what lies behind the virtual
80 performance. Contesting gender norms through technology
81 and Japanese theatre. *Electronic Journal of Contemporary
82 Japanese Studies*, 22(2), 1-22.
- 83 Broom, T. W., Chavez, R. S., & Wagner, D. D. (2021). Becoming
84 the King in the North: identification with fictional characters
85 is associated with greater self-other neural overlap. *Social
86 Cognitive and Affective Neuroscience*, 16(6), 541-551.
- 87 Brown, S., Cockett, P., & Yuan, Y. (2019). The neuroscience of
88 Romeo and Juliet: an fMRI study of acting. *Royal Society
89 Open Science*, 6(3), 181908.
- 90 Bruckert, L., Bestelmeyer, P., Latinus, M., Rouger, J., Charest, I.,
91 Rousselet, G. A., Kawahara, H., & Belin, P. (2010). Vocal
92 attractiveness increases by averaging. *Current Biology*,
93 20(2), 116-120.
- 94 Chiba, Y., & Ito, A. (2024). Speaker intimacy estimation in chat-

- 1 talks based on verbal and non-verbal information. *IEEE*
- 2 *Access*, 12, 184592-184606.
- 3 Cohn, M., Bhandodkar, G., Sangani, R. B., Predeck, K., & Zellou,
- 4 G. (2024). Do people mirror emotion differently with a
- 5 human or TTS voice? Comparing listener ratings and word
- 6 embeddings. *Extended Abstracts of the CHI Conference on*
- 7 *Human Factors in Computing Systems*, 1-10.
- 8 Crochiquia, A., Eriksson, A., Fontes, M. A., & Madureira, S.
- 9 (2020). A phonetic study of Zootopia characters' voices in
- 10 Brazilian Portuguese dubbing: the role of stereotypes.
- 11 *DELTA: Documentação de Estudos em Lingüística Teórica*
- 12 *e Aplicada*, 36(3), 1-46.
- 13 Daryadar, M., & Raghbi, M. (2015). The effect of listening to
- 14 recordings of one's voice on attentional bias and auditory
- 15 verbal learning. *International Journal of Psychological*
- 16 *Studies*, 7(2), 155-163.
- 17 藤津 亮太 (2017). 声優語:アニメに命を吹き込むプロフェ
- 18 ヂショナル 一迅社
- 19 藤津 亮太 (2018). 声優論:通史的, 実証的一考察 小川 昌
- 20 宏・須川 亜紀子 (編著) アニメ研究入門【応用編】:ア
- 21 ニメを究める 11 つのコツ (pp. 93-117) 現代書館
- 22 藤津 亮太 (2019). プロフェッショナル 13 人が語るわたし
- 23 の声優道 河出書房新社
- 24 福田 航希・阪井 瞭介・松下 嶺佑・國見 友亮・高道 慎之
- 25 介 (2025). ELEVATE: 学習者自身の自己聴取音声で聴
- 26 く講義システム インタラクシオン2025 論文集, 1B-21,
- 27 314-319.
- 28 布山 美慕・西郷 甲矢人 (2022). 解釈の不定性の価値と量
- 29 子認知による文章解釈研究の展望 認知科学, 29(1),
- 30 100-119.
- 31 Gallagher, S. (2000). Philosophical conceptions of the self:
- 32 implications for cognitive science. *Trends in Cognitive*
- 33 *Sciences*, 4(1), 14-21.
- 34 Goldstein, T. R., & Filipe, A. (2018). The interpreted mind:
- 35 Understanding acting. *Review of General Psychology*, 22(2),
- 36 220-229.
- 37 Grawe, G. (2018). 日本文化における「声」 立命館言語文化
- 38 研究, 29, 155-173.
- 39 原 雄太郎・伊藤 克亘 (2009). 声優の発話の音響特徴量分
- 40 析及び確率モデルの作成 情報科学技術フォーラム講
- 41 演論文集, 8(2), 369-372.
- 42 Hatada, Y., Barbareschi, G., Takeuchi, K., Kato, H., Yoshifuji, K.,
- 43 Minamizawa, K., & Narumi, T. (2024). People with
- 44 disabilities redefining identity through robotic and virtual
- 45 avatars: A case study in avatar robot cafe. *Proceedings of the*
- 46 *CHI Conference on Human Factors in Computing Systems*,
- 47 61, 1-13.
- 48 林 大輔 (2019). 声優のキャラクター演技音声を用いた音
- 49 声知覚に関する実験研究 愛知淑徳大学論集一人間
- 50 情報学部篇, 9, 49-62.
- 51 林 大輔・森勢 将雅 (2025). 女性声優の演技音声における
- 52 年齢・性別の表現と関連する音響特徴量 Jxiv.
- 53 林 大輔・大杉 尚之 (2021). アニメにおける声質に対する
- 54 印象の調査: 日常場面およびドラマとの比較 PsyArXiv.
- 55 Holzman, P. S., & Rousey, C. (1966). The voice as a percept.
- 56 *Journal of Personality and Social Psychology*, 4(1), 79-86.
- 57 本多 清志 (2018). 実験音声科学: 音声事象の成立過程を探
- 58 る コロナ社
- 59 本間 裕之 (2024). 「身体」と「魂」としての VTuber 岡本
- 60 健・山野 弘樹・吉川 慧 (編著) VTuber 学 (pp. 323-337)
- 61 岩波書店
- 62 Hosaka, T., Kimura, M., & Yotsumoto, Y. (2021). Neural
- 63 representations of own-voice in the human auditory cortex.
- 64 *Scientific Reports*, 11(1), 591.
- 65 Hughes, S. M., & Harrison, M. A. (2013). I like my voice better:
- 66 Self-enhancement bias in perceptions of voice attractiveness.
- 67 *Perception*, 42(9), 941-949.
- 68 Iizuka, T., & Mori, H. (2022). How does a spontaneously
- 69 speaking conversational agent affect user behavior? *IEEE*
- 70 *Access*, 10, 111042-111051.
- 71 稲岡 大志 (2016). 堀江由衣をめぐる試論: 音声・キャラク
- 72 ター・同一性 フィルカル, 1(2), 112-140.
- 73 井上 晴菜 (2022). 言語的符号化が標的音声の話者同定に
- 74 与える影響 心理学研究, 93(4), 320-329.
- 75 石田 美紀 (2020). アニメと声優のメディア史: なぜ女性が
- 76 少年を演じるのか 青弓社
- 77 Ishi, C.T., Utsugi, A., & Ota, I. (2023). Voice types and voice
- 78 quality in Japanese anime. *Proceedings of the 20th*
- 79 *International Congress of Phonetic Sciences*, 3632-3636.
- 80 石井 沙季・伊藤 克亘 (2019). キャラクター音声のステレ
- 81 オタイプ識別のための音響分析 情報処理学会第 81 回
- 82 全国大会講演論文集, 4, 695-696.
- 83 ジン リーフアン (2010). 日本のアニメーションにおける
- 84 音声の機能: 「GHOST IN THE SHELL/攻殻機動隊」「新
- 85 世紀エヴァンゲリオン」「蒼穹のファフナー」を中心に
- 86 北海道大学大学院文学研究科研究論集, 10, 235-251.
- 87 Jones, B. C., Feinberg, D. R., DeBruine, L. M., Little, A. C., &
- 88 Vukovic, J. (2010). A domain-specific opposite-sex bias in
- 89 human preferences for manipulated voice pitch. *Animal*
- 90 *Behaviour*, 79(1), 57-62.
- 91 Jürgens, R., Hammerschmidt, K., & Fischer, J. (2011). Authentic
- 92 and play-acted vocal emotion expressions reveal acoustic
- 93 differences. *Frontiers in Psychology*, 2(180), 1-11.
- 94 Kao, D., Ratan, R., Mousas, C., & Magana, A. J. (2021). The
- 95 effects of a self-similar avatar voice in educational games.
- 96 *Proceedings of the ACM on Human-Computer Interaction*,
- 97 5(CHI PLAY), 1-28.
- 98 Kawahara, S. (2016). The prosodic features of the "moe" and
- 99 "tsun" voices. *Journal of the Phonetic Society of Japan*,
- 100 20(2), 102-110.
- 101 川村 寛文 (2024). 聖なるもの, 情動, プラットフォーム:
- 102 声優/キャラ・ライブコンサートにみるリアリティの複
- 103 数性 須川 亜紀子 (編) 2.5 次元学入門 (pp. 81-113) 青
- 104 土社
- 105 木戸 博・粕谷 英樹 (1999). 通常発話の声質に関連した日
- 106 常表現語の抽出 日本音響学会誌, 55(6), 405-411.

- 1 Kimura, M., & Yotsumoto, Y. (2018). Auditory traits of “own
2 voice.” *PLoS ONE*, 13(6), e0199443.
- 3 Kitamura, T., Honda, K., & Takemoto, H. (2005). Individual
4 variation of the hypopharyngeal cavities and its acoustic
5 effects. *Acoustical Science and Technology*, 26(1), 16-26.
- 6 Krauss, R. M., Freyberg, R., & Morsella, E. (2002). Inferring
7 speakers’ physical attributes from their voices. *Journal of*
8 *Experimental Social Psychology*, 38, 618-625.
- 9 久保(川合) 南海子 (2022). 「推し」の科学：プロジェクシ
10 ョン・サイエンスとは何か 集英社新書
- 11 國見 友亮・畑田 裕二・木村 健太・鳴海 拓志・持丸 正明
12 (2024). 声質変化に伴う自己意識の変化についてのア
13 ンケート調査 第29回日本バーチャルリアリティ学会
14 大会論文集, 2D1-08.
- 15 Kunimi, Y., Kimura, K., Matsumoto, K., Takamichi, S., Narumi,
16 T., & Mochimaru, M. (2024). Character-voice embodiment
17 impacts on the cognitive task performance with the voice
18 ownership illusion. *International Conference on Artificial*
19 *Reality and Telexistence Eurographics Symposium on*
20 *Virtual Environments 2024*, 1-10.
- 21 倉井 龍太郎・平木 剛史 (2023). ソーシャル VR プラット
22 フォームにおけるエージェント API の提案 第28回日
23 本バーチャルリアリティ学会大会論文集, 3D1-08.
- 24 倉田 将希・高道 慎之介・佐伯 高明・荒川 陸・齋藤 佑樹・
25 樋口 啓太・猿渡 洋 (2021). リアルタイム DNN 音声
26 変換フィードバックによるキャラクター性の獲得手法
27 研究報告音声言語情報処理(SLP), 2021(31), 1-6.
- 28 Kuriki, S., Tamura, Y., Igarashi, M., Kato, N., & Nakano, T.
29 (2016). Similar impressions of humanness for human and
30 artificial singing voices in autism spectrum disorders.
31 *Cognition*, 153, 1-5.
- 32 栗田 茂二郎 (1988). 声帯の成長, 発達と老化: とくに層構
33 造の加齢的变化 音声言語医学, 29(2), 185-193.
- 34 黒崎 想 (2018). 縫い付けられた声 ユリイカ, 50(9), 188-
35 195.
- 36 Latinus, M., & Belin, P. (2011). Anti-voice adaptation suggests
37 prototype-based coding of voice identity. *Frontiers in*
38 *Psychology*, 2(175), 1-12.
- 39 Laukka, P., Neiberg, D., Forsell, M., Karlsson, I., & Elenius, K.
40 (2011). Expression of affect in spontaneous speech: Acoustic
41 correlates and automatic detection of irritation and
42 resignation. *Computer Speech and Language*, 25(1), 84-104.
- 43 Lu, Z., Shen, C., Li, J., Shen, H., & Wigdor, D. (2021). More
44 kawaii than a real-person live streamer: Understanding how
45 the otaku community engages with and perceives virtual
46 YouTubers. *Proceedings of the 2021 CHI Conference on*
47 *Human Factors in Computing Systems*, 1-14.
- 48 Lundh, L. G., Berg, B., Johansson, H., Nilsson, L. K., Sandberg,
49 J., & Segerstedt, A. (2002). Social anxiety is associated with
50 a negatively distorted perception of one's own voice.
51 *Cognitive Behaviour Therapy*, 31(1), 25-30.
- 52 丸島 歩 (2020a). 女性声優による役柄の性別の異なる音声
53 の音響的特徴-基本周波数に着目して 大阪経済法科大
54 学論集, 115, 23-33.
- 55 丸島 歩 (2020b). 女性声優の演技音声にあらわれるジェン
56 ダーの表現: 母音フォルマントに着目して 年報新人
57 文学, 17, 165-139.
- 58 増野 亜子 (2014). 声の世界を旅する 音楽之友社
- 59 松本 大輝 (2017). その歌は緑の髪をしている: ボーカロイ
60 ドとメイクビリーブ フィルカル, 2(2), 96-142.
- 61 松本 大輝 (2024). フィクショナル・キャラクターとしての
62 VTuber 岡本 健・山野 弘樹・吉川 慧 (編著) VTuber 学
63 (pp. 307-322) 岩波書店
- 64 Matsunaga, Y., Saeki, T., Takamichi, S., & Saruwatari, H. (2022).
65 *Improving robustness of spontaneous speech synthesis with*
66 *linguistic speech regularization and pseudo-filled-pause*
67 *insertion*. arXiv.
- 68 McAleer, P., Todorov, A., & Belin, P. (2014). How do you say
69 ‘Hello’? Personality impressions from brief novel voices.
70 *PLoS ONE*, 9(3), e90779.
- 71 三原 鉄也 (2022). メディア芸術に関する行政施策の展開
72 とデジタルアーカイブ デジタルアーカイブ学会誌,
73 6(1), 15-19.
- 74 三谷 雅純 (2014). 生涯学習施設の館内放送はどうあるべ
75 きか: 聴覚実験による肉声と人工合成音声の聞きやす
76 きの比較 人と自然, 25, 63-74.
- 77 三井 康行 (2010). 音声合成の利用シーンと要求される品
78 質との関係 研究報告音声言語情報処理(SLP), 2010(8),
79 1-4.
- 80 Mitsui, K., Hono, Y., & Sawada, K. (2023). *Towards human-like*
81 *spoken dialogue generation between ai agents from written*
82 *dialogue*. arXiv.
- 83 宮島 崇浩・菊池 英明・白井 克彦・大川 茂樹 (2013). 演
84 技指示の工夫が与える音声表現への影響: 表現豊かな
85 演技音声表現の獲得を目指して 音声研究, 17(3), 10-23.
- 86 宮本 直美 (2017). 嗜好対象としての歌声: クラシック歌唱
87 からポピュラー歌唱へ 嗜好品文化研究, 2017(2), 26-37.
- 88 モクタリ 明子 (2008). 自然発話にみられる人物像に応じ
89 た個人内音声バリエーション 神戸大学総合人間科学
90 研究科博士学位論文
- 91 森 大毅・前川 喜久雄・粕谷 英樹 (2014). 音声は何を伝え
92 ているか: 感情・パラ言語情報・個人性の音声科学 コ
93 ロナ社
- 94 村瀬 ひろみ・須川 亜紀子 (2014). ジェンダー論 アニメか
95 ら読み解くジェンダー: 「力」と「美」を超えて 小川
96 昌宏・須川 亜紀子 (編著) 増補改訂版アニメ研究入
97 門: アニメを究める9つのツボ (pp. 76-95) 現代書館
- 98 永田 大輔 (2024). アイドル声優はなぜ問題となるのか: メ
99 ディア史の中の「キャリア」 メディウム, 5, 19-34.
- 100 中川 優奈・田中 章浩 (2021). 自分と他人の声の境界は変
101 化するか 信学技報, 121(177), 25-30.
- 102 難波 優輝 (2018). バーチャル YouTuber の三つの身体: パ
103 ーソン, ペルソナ, キャラクタ ユリイカ, 50(9), 117-125.
- 104 鳴海 拓志 (2019). ゴーストエンジニアリング: 身体変容に
105 よる認知拡張の活用に向けて 認知科学, 26(1), 14-29.
- 106 成瀬 加菜・吉田 成朗・高道 慎之介・鳴海 拓志・谷川 智

- 1 洋・廣瀬 通孝 (2019). 自信声フィードバックによる緊
2 張緩和手法の提案: クラウドソーシングを利用した自
3 信声加工パラメータの推定 第24回日本バーチャルリ
4 アリティ学会大会論文集, 1B-05.
- 5 Nguyen, T. A., Kharitonov, E., Copet, J., Adi, Y., Hsu, W.-N.,
6 Elkahky, A., Tomasello, P., Algayres, R., Sago, B.,
7 Mohamed, A., & Dupoux, E. (2023). Generative spoken
8 dialogue language modeling. *Transactions of the*
9 *Association for Computational Linguistics*, 11, 250-266.
- 10 日本音響学会(編)(1996). 音のなんでも小事典: 脳が音を聴
11 くしくみから超音波顕微鏡まで 講談社
- 12 日本音響学会(編)(2003). 新版 音響用語辞典 コロナ社
- 13 Ogawa, N., Baba, J., & Nakanishi, J. (2024). Investigating effect
14 of altered auditory feedback on self-representation,
15 subjective operator experience, and task performance in
16 teleoperation of a social robot. *Proceedings of the 2024 CHI*
17 *Conference on Human Factors in Computing Systems*, 1-18.
- 18 Ohata, R., Asai, T., Imaizumi, S., & Imamizu, H. (2022). I hear
19 my voice; therefore I spoke: The sense of agency over
20 speech is enhanced by hearing one's own voice.
21 *Psychological science*, 33(8), 1226-1239.
- 22 Ostrega, J., Shiramizu, V., Lee, A. J., Jones, B. C., & Feinberg, D.
23 R. (2024). No evidence that averaging voices influences
24 attractiveness. *Scientific Reports*, 14(1), 10488.
- 25 皇牙 サキ (2018). 「声」という商品のパッケージとしての
26 VTuber ユリイカ, 50(9), 64-65.
- 27 Peng, Z., Wang, Y., Meng, L., Liu, H., & Hu, Z. (2019). One's
28 own and similar voices are more attractive than other voices.
29 *Australian Journal of Psychology*, 71(3), 212-222.
- 30 Qi, T., Zheng, W., Lu, C., Zong, Y., & Lian, H. (2024). Pavits:
31 Exploring prosody-aware vits for end-to-end emotional
32 voice conversion. *IEEE International Conference on*
33 *Acoustics, Speech and Signal Processing*, 12697-12701.
- 34 Saeki, T., Takamichi, S., Nakamura, T., Tanji, N., & Saruwatari,
35 H. (2023). Selfmaster: Self-supervised speech restoration
36 for historical audio resources. *IEEE Access*, 11, 144831-
37 144843.
- 38 Saito, Y., Nishimura, Y., Takamichi, S., Tachibana, K., &
39 Saruwatari, H. (2022). *STUDIES: Corpus of Japanese*
40 *empathetic dialogue speech towards friendly voice agent*.
41 arXiv.
- 42 Samo, A., & Highhouse, S. (2023). Artificial intelligence and art:
43 Identifying the aesthetic judgment factors that distinguish
44 human-and machine-generated artwork. *Psychology of*
45 *Aesthetics, Creativity, and the Arts*, Advance online
46 publication.
- 47 佐藤 茉奈花 (2023). 同一声優による異なる性格を持つキ
48 ャラクターの演技音声の分析 社会言語科学会第47回
49 大会発表論文集, 2-3, 29-32.
- 50 さやわか (2015). キャラの思考法: 現代文化論のアップグ
51 レード 青土社
- 52 Scherer, K. R. (1972). Judging personality from voice: A cross-
53 cultural approach to an old issue in interpersonal perception.
54 *Journal of Personality*, 40(2), 191-210.
- 55 Scherer, K. R. (1978). Personality inference from voice quality:
56 The loud voice of extroversion. *European Journal of Social*
57 *Psychology*, 8(4), 467-487.
- 58 Scherer, K. R. (2003). Vocal communication of emotion: A review
59 of research paradigms. *Speech Communication*, 40(1-2),
60 227-256.
- 61 Schweinberger, S. R., Casper, C., Hauthal, N., Kaufmann, J. M.,
62 Kawahara, H., Kloth, N., Robertson, D. M. C., Simpson, A.
63 P., & Zäske, R. (2008). Auditory adaptation in voice
64 perception. *Current Biology*, 18(9), 684-688.
- 65 Schweinberger, S. R., Kawahara, H., Simpson, A. P., Skuk, V. G.,
66 & Zaske, R. (2014). Speaker perception. *WIREs Cognitive*
67 *Science*, 5, 15-25.
- 68 関根 麻里恵 (2024). メタVTuber コンテンツの表象文化研
69 究: 「匿名性」「有名性」「声」「ジェンダー」から考え
70 る 岡本 健・山野 弘樹・吉川 慧 (編著)VTuber 学 (pp.
71 179-195) 岩波書店
- 72 柴 那典 (2014). 初音ミクはなぜ世界を変えたのか? 太田
73 出版
- 74 嶋田 総太郎 (2019). 脳のなかの自己と他者: 身体性・社会
75 性の認知脳科学と哲学 共立出版
- 76 標葉 隆馬・見上 公一 (編)(2024). 入門 科学技術と社会
77 ナカニシヤ出版
- 78 シノハラ ユウキ (2017). メディアを跨ぐヴィヴィッドな
79 想像: 『Tokyo 7th シスターズ』における「跳ぶよ」と
80 いうセリフの事例から フィルカル, 2(2), 60-95.
- 81 篠崎 大河 (2024). 実在する配信者としての VTuber 岡本
82 健・山野 弘樹・吉川 慧 (編著)VTuber 学 (pp. 271-288)
83 岩波書店
- 84 Sisman, B., Yamagishi, J., King, S., & Li, H. (2021). An overview
85 of voice conversion and its challenges: From statistical
86 modeling to deep learning. *IEEE/ACM Transactions on*
87 *Audio, Speech, and Language Processing*, 29, 132-157.
- 88 Starr, R. L. (2015). Sweet voice: The role of voice quality in a
89 Japanese feminine style. *Language in Society*, 44(1), 1-34.
- 90 Suda, H., Watanabe, A., & Takamichi, S. (2024). Who finds this
91 voice attractive? A large-scale experiment using in-the-wild
92 data. *Proceedings of Interspeech 2024*, 3165-3169.
- 93 須川 亜紀子 (2021). 2.5次元文化論: 舞台・キャラクター・
94 ファンダム 青弓社
- 95 鈴木 真吾 (2014). サウンド/ヴォイス研究 アニメを奏で
96 る3つの音: アニメにとって音とは何か 小川 昌宏・
97 須川 亜紀子 (編著) 増補改訂版アニメ研究入門: アニ
98 メを究める9つのツボ (pp. 96-119) 現代書館
- 99 高道 慎之介 (2021). 音声アバターを選ぶ時代: ボイスチェ
100 ンジャー技術の動向 電気学会誌, 141(2), 93-96.
- 101 Tan, X., Chen, J., Liu, H., Cong, J., Zhang, C., Liu, Y., Wang, X.,
102 Leng, Y., Yi, Y., He, L., Soong, F., Qin, T., Zhao, S., & Liu,
103 T.-Y. (2024). NaturalSpeech: End-to-end text-to-speech
104 synthesis with human-level quality. *IEEE Transactions on*
105 *Pattern Analysis and Machine Intelligence*, 46(6), 4234-
106 4245.

- 1 田中 彰吾 (編著)(2023). 自己の科学は可能か：心身脳問題
2 として考える 新曜社
- 3 程 斯 (2023). アニメ・キャラクターとアニメ声優を同一視
4 する現象の原理的基盤をめぐって *Phantastopia*, 2, 122-
5 139.
- 6 Teshigawara, M. (2004). Vocally expressed emotions and
7 stereotypes in Japanese animation: Voice qualities of the bad
8 guys compared to those of the good guys. *Journal of the*
9 *Phonetic Society of Japan*, 8(1), 60-76.
- 10 勅使河原 三保子 (2019). 声に関するステレオタイプ of 解
11 明に向けて：音声に基づく人物像の知覚の 3 次元モ
12 デル 駒澤大学外国語論集, 27, 1-19.
- 13 勅使河原 三保子・伊藤 克亘・武田 一哉 (2005). 日本のア
14 ニメの音声に表された感情と性格：声のステレオタイ
15 プの音声学的研究 電子情報通信学会技術研究報告,
16 105(291), 39-44.
- 17 Titze, I. R. (1989). Physiologic and acoustic differences between
18 male and female voices. *Journal of the Acoustical Society of*
19 *America*, 85(4), 1699-1707.
- 20 Uchida, T., Takahashi, H., Ban, M., Shimaya, J., Minato, T.,
21 Ogawa, K., Yoshikawa, Y., & Ishiguro, H. (2020). Japanese
22 young women did not discriminate between robots and
23 humans as listeners for their self-disclosure-pilot study.
24 *Multimodal Technologies and Interaction*, 4(35), 1-16.
- 25 内田 照久 (2000). 音声の発話速度の制御がピッチ感及び
26 話者の性格印象に与える影響 日本音響学会誌, 56(6),
27 396-405.
- 28 内田 照久 (2002). 音声の発話速度が話者の性格印象に与
29 える影響 心理学研究, 73(2), 131-139.
- 30 内田 照久 (2005a). 音声の発話速度と休止時間が話者の性
31 格印象と自然なわかりやすさに与える影響 教育心理
32 学研究, 53(1), 1-13.
- 33 内田 照久 (2005b). 音声の中の抑揚の大きさと変化パターン
34 が話者の性格印象に与える影響 心理学研究, 76(4),
35 382-390.
- 36 内田 照久 (2006). 未知のイントネーションから想起され
37 る話者の性格印象と方言地域の特徴 音声研究, 10(3),
38 29-42.
- 39 内田 照久 (2009). 音声の韻律的特徴と話者のパーソナリ
40 ティ印象の関係性 音声研究, 13(1), 17-28.
- 41 内田 照久 (2011). 音声の中の母音の明瞭性が話者の性格印
42 象と話し方の評価に与える影響 心理学研究, 82(5),
43 433-441.
- 44 内田 照久・森勢 将雅 (2020). 声のピッチ感の錯覚と疑似
45 歌声・疑似ささやき声による検討 情報処理学会論文
46 誌, 61(4), 807-816.
- 47 内田 照久・中畝 菜穂子 (2004). 声の高さと発話速度が話
48 者の性格印象に与える影響 心理学研究, 75(5), 397-406.
- 49 Ueda, A., Takahashi, H., Yoshikawa, Y., Ishiguro, H., & Nomura,
50 H. (2021). Do robots facilitate life review narratives of older
51 adults? A preliminary study. *Gerontechnology*, 20(2), 1-12.
- 52 Utsugi, A., Wang, H., & Ota, I. (2019). A voice quality analysis
53 of Japanese anime. *Proceedings of the 19th International*
54 *Congress of the Phonetic Sciences*, 1853-1857.
- 55 Van Lancker, D., Kreiman, J., & Emmorey, K. (1985). Familiar
56 voice recognition: patterns and parameters Part I:
57 Recognition of backward voices. *Journal of Phonetics*, 13,
58 19-38.
- 59 Vorperian, H. K., Wang, S., Chung, M. K., Schimek, E. M.,
60 Durtschi, R. B., Kent, R. D., Ziegert, A. J., & Gentry, L. R.
61 (2009). Anatomic development of the oral and pharyngeal
62 portions of the vocal tract: An imaging study. *Journal of the*
63 *Acoustical Society of America*, 125(3), 1666-1678.
- 64 Vorperian, H. K., Wang, S., Schimek, E. M., Durtschi, R. B., Kent,
65 R. D., Gentry, L. R., & Chung, M. K. (2011). Developmental
66 sexual dimorphism of the oral and pharyngeal portions of the
67 vocal tract: An imaging study. *Journal of Speech, Language,*
68 *and Hearing Research*, 54(4), 995-1010.
- 69 Walton, K. L. (1990). Mimesis as make-believe: On the
70 foundations of the representational arts. Harvard University
71 Press.
- 72 (ウォルトン・ケンダル 田村 均 (訳) (2016). フィク
73 ションとは何か：ごっこ遊びと芸術 名古屋大学出版
74 会)
- 75 Weiss, B., Trouvain, J., Barkat-Defradas, M., & Ohala, J. J. (Ed.)
76 (2021). *Voice attractiveness: Studies on sexy, likable, and*
77 *charismatic speakers*. Springer Singapore.
- 78 Williams, C. E., & Stevens, K. N. (1972). Emotions and speech:
79 Some acoustical correlates. *Journal of the Acoustical Society*
80 *of America*, 52(4B), 1238-1250.
- 81 Xin, D., Takamichi, S., Morimatsu, A., & Saruwatari, H. (2023).
82 Laughter synthesis using pseudo phonetic tokens with a
83 large-scale in-the-wild laughter corpus. *Proceedings of*
84 *Interspeech 2023*, 17-21.
- 85 山田 真也・伊藤 敏彦・荒木 健治 (2005). 対話相手の音声
86 の品質を考慮した対話状況での言語的・音響的特徴の
87 分析および様々な観点からの考察 情報処理学会研究
88 報告音声言語情報処理(SLP), 2005(127), 67-72.
- 89 山中 涼雅・大佐 健人・藤原 朱里・耿 浩彭・齋藤 大輔・
90 峯松 信明・井上 雄介 (2025). 学習者本人の自己聴取
91 音の声質で合成されたモデル音声を用いた発音学習
92 とその効果 2025 年春季日本音響学会研究発表会講演
93 論文集, 1165-1168.
- 94 山野 弘樹 (2024). VTuber の哲学 春秋社
- 95 山野 弘樹 (2025). 「VTuber 文化」を哲学する第 29 回
96 VTuber スタイル, 2025(6), 78.
- 97 Yanagida, H., Ijima, Y., & Tawara, N. (2023). Influence of
98 personal traits on impressions of one's own voice.
99 *Proceedings of Interspeech 2023*, 5212-5216.
- 100 横森 文哉・二宮 大和・森勢 将雅・田中 章浩・小澤 賢司
101 (2016). 好感度評価の性差に着目した女性発話の音響
102 特徴量分析 日本感性工学会論文誌, 15(7), 721 -729.
- 103