

CURRICULUM VITAE

Keiichi Tokuda

Department of Computer Science
Nagoya Institute of Technology
Gokiso-cho, Showa-ku
Nagoya, 466-8555 JAPAN

June 9, 2026

Email: tokuda [at] nitech.ac.jp

WWW: <https://www.sp.nitech.ac.jp/~tokuda/>

1 PROFESSIONAL EXPERIENCE

April 2026–present

Visiting Professor/Professor Emeritus, Department of Computer Science, Nagoya Institute of Technology, Japan

April 2004–2026

Professor, Department of Computer Science, Nagoya Institute of Technology, Japan

April 1996–March 2004

Associate Professor, Department of Computer Science, Nagoya Institute of Technology, Japan

April 1989–March 1996

Research Associate, Department of Electrical and Electronic Engineering, Tokyo Institute of Technology, Japan

January 2012–2017

Honorary Professor, University of Edinburgh, U.K.

April 2014–March 2015

Visiting Researcher, Google, U.K.

April 2006–March 2013

Visiting Researcher, National Institute of Information and Communications Technology (NiCT), Japan

April 2002–March 2006

Visiting Researcher, ATR Spoken Language Communication Laboratories, Japan

July 2001–April 2002

Visiting Researcher, Language Technologies Institute Carnegie Mellon University

April 2001–September 2001

Invited Lecturer, Nagoya University, Japan

August 2000–July 2001

Visiting Researcher, ATR Spoken Language Translation Laboratories, Japan

November 1997–March 1998

Invited Lecturer, Tokyo Institute of Technology, Japan

2 UNIVERSITY DEGREES

1989

Dr.Eng in information processing, Tokyo Institute of Technology, Japan

1986

M.E in information processing, Tokyo Institute of Technology, Japan

1984

B.E in electrical and electronic engineering, Nagoya Institute of Technology, Japan

3 AWARDS

2024 Asia-Pacific Artificial Intelligence Association (AAIA) Fellow

2024 IEEE James L. Flanagan Speech and Audio Processing Award

2020 The Medal with Purple Ribbon, Government of Japan

2019 ISCA Best Paper Award, 6th ISCA Speech Synthesis Workshop

2019 ISCA Medal for Scientific Achievement

2015 Achievement Award from the Institute of Electronics, Information and Communication Engineers

2014 IEEE Fellow

2013 ISCA Fellow

2013 IPSJ Kiyasu Special Industrial Achievement Award (Information Processing Society of Japan)

2013 The 2013 EURASIP-ISCA Best Paper Award (Speech Communication Journal)

2012 Prizes for Science and Technology (Research Category), The Commendation for Science and Technology by the Minister of Education, Culture, Sports, Science and Technology

2008 Information and Systems Society Distinguished Achievement Award from the Institute of Electronics, Information and Communication Engineers

- 2008** Information and Systems Society Excellent Paper Award from the Institute of Electronics, Information and Communication Engineers
- 2008** TELECOM System Technology Prize from the Telecommunications Advancement Foundation
- 2001** Best Paper Award from the Institute of Electronics, Information and Communication Engineers
- 2001** Inose Award (the highest award) from the Institute of Electronics, Information and Communication Engineers
- 2001** TELECOM System Technology Prize from the Telecommunications Advancement Foundation

4 PUBLICATION

BOOK CHAPTERS

1. Keiichi Tokuda, Akinobu Lee, Yoshihiko Nankaku, Keiichiro Oura, Kei Hashimoto, Daisuke Yamamoto, Ichi Takumi, Takahiro Uchiya, Shuhei Tsutsumi, Steve Renals, and Junichi Yamagishi, “User Generated Dialogue Systems: uDialogue,” Human-Harmonized Information Technology, Volume 2, Springer, pp.77–114, May 2017. (ISBN 978-4-431-56533-8)
2. Heiga Zen, Keiichi Tokuda, “7.3 The HMM-based speech synthesis system (HTS),” in Computer processing of Asian spoken languages, Editors: Shuichi Itahashi, Chiu-yu Tseng, Consideration Books, Los Angeles, March 2010. (ISBN 978-0-935047-72-1)
3. Keiichi Tokuda, “4.3 A unified approach to prosody control using HMM-based speech synthesis,” Prosody and Spoken Language Processing, Keikichi Hirose (Ed.), Maruzen, pp.118–127, Jan. 2006 (in Japanese). (ISBN 978-4-621-07674-3)
4. Keiichi Tokuda, Heiga Zen, Alan W. Black, “An HMM-Based Approach to Multilingual Speech Synthesis (Chapter 7),” Text-to-Speech Synthesis: New Paradigms and Advances, Shrikanth Narayanan, Abeer Alwan (Eds.), Prentice Hall, pp.135–153, Aug. 2004. (ISBN 978-0131456617)
5. Shin-ichi Kawamoto, Hiroshi Shimodaira, Tsuneo Nitta, Takuya Nishimoto, Satoshi Nakamura, Katsunobu Itou, Shigeo Morishima, Tatsuo Yotsukura, Atsuhiko Kai, Akinobu Lee, Yoichi Yamashita, Takao Kobayashi, Keiichi Tokuda, Keikichi Hirose, Nobuaki Minematsu, Atsushi Yamada, Yasuharu Den, Takehito Utsuro, Shigeki Sagayama, “Galatea: Open-source Software for Developing Anthropomorphic Spoken Dialog Agents,” Life-Like Characters: Tools, Affective Functions, and Applications, Series: Cognitive Technologies, Helmut Prendinger, Mitsuru Ishizuka (Eds.), Springer-Verlag, pp.187–211, 2004. (ISBN 978-3540008675)
6. Takao Kobayashi, Keiichi Tokuda, “14.3 Speech Spectral Estimation,” in Handbook of Spectral Analysis, Mikio Hino (Ed.), Asakura, pp.481–492, 2004 (in Japanese). (ISBN 978-4-254-20108-6)

INVITED REVIEW AND TUTORIAL

1. Keiichi Tokuda, "Recent Trends in Speech Synthesis Research," IEICE Information and Systems Society Journal, vol.21, no.4, pp.10–11, April, 2017 (in Japanese).
2. Keiichi Tokuda, Yoshihiko Nankaku, Tomoki Toda, Heiga Zen, Junichi Yamagishi, Keiichi Oura, "Speech Synthesis Based on Hidden Markov Models," Proceedings of the IEEE, vol.101, no.5, pp.1234–1252, May 2013.
3. Keiichi Tokuda, "Recent advances in statistical parametric speech synthesis," the Journal of the Acoustical Society of Japan, vol.67, no.1, pp.17–22, January 2011 (in Japanese).
4. Keiichi Oura, Heiga Zen, Shinji Sako, Keiichi Tokuda, "Construction of Speech Synthesis Systems using HTS," Journal of Human Interface Society : human interface, vol.12, no.1, pp.35–40, February 2010 (in Japanese).
5. Hisashi Kawai, Keiichi Tokuda, "Multi-language Speech Synthesis," The Journal of The Institute of Electrical Engineers of Japan, vol.130, no.1, pp.16–19, January 2010 (in Japanese).
6. Kiyohiro Shikano, Kazuya Takeda, Tatsuya Kawahara, Hideki Kawahara, Hiroshi Saruwatari, Keiichi Tokuda, Akinobu Lee, Hiromichi Kawanami, Ryuichi Nishimura, Randy GOMEZ, Tomoki Toda, Takanobu Nishiura, Toru Takahashi, Hideki Banno, Heiga Zen, "E-Society Software Development Project for Speech Recognition and Synthesis," Journal of IEICE, vol.92, no.6, June 2009 (in Japanese).
7. Heiga Zen, Keiichi Tokuda, "TechWare: HMM-Based Speech Synthesis Resources," IEEE Signal Processing Magazine, July 2009 (in Japanese).
8. Kiyohiro Shikano, Tatsuya Kawahara, Hiroshi Saruwatari, Kazuya Takeda, Hideki Kawahara, Keiichi Tokuda, Takanobu Nishiura, Akinobu Lee, "Advanced development of fundamental software through tight collaboration of academia and industry: Human friendly speech interface," IPSJ Magazine, vol.49, no.11, pp.1297–1301, November 2008 (in Japanese).
9. Keiichi Tokuda, Alan W. Black, "Speech synthesis research in a new age of cooperation and competition: The Blizzard Challenge," the Journal of the Acoustical Society of Japan, vol.62, no.6, pp.466–472, June 2006 (in Japanese).
10. Keiichi Tokuda, "Speech Recognition and Speech Synthesis based on Hidden Markov Models," IPSJ Magazine, vol.45, no.10, (in Japanese).
11. Takao Kobayashi, Keiichi Tokuda, "Technology Trends in Corpus-based Speech Synthesis [IV]: HMM-Based Speech Synthesis," The Journal of IEICE, vol.87, no.4, April 2004 (in Japanese).

INVITED PAPER

1. Keiichi Oura, Daisuke Yamamoto, Ichi Takumi, Akinobu Lee, Keiichi Tokuda, "On-Campus, User-Participatable, and Voice-Interactive Digital Signage," Journal of The Japanese Society for Artificial Intelligence, vol.28, no.1, pp.60–67, January 2013.

2. Keiichi Tokuda, Takashi Masuko, Noboru Miyazaki, Takao Kobayashi, “Multi-space probability distribution HMM,” *IEICE Trans. Information and Systems*, vol.E85-D, no.3, pp.455–464, Mar. 2002.

REFEREED JOURNAL PAPERS

1. Yoshihiko Nankaku, Takato Fujimoto, Takenori Yoshimura, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, and Keiichi Tokuda, “Deep Hidden Semi-Markov Model-Based Speech Synthesis,” in *IEEE Access*, vol. 14, pp. 58495-58514, 2026.
2. Takato Fujimoto, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, “V2Coder: A Non-Autoregressive Vocoder Based on Hierarchical Variational Autoencoders,” in *IEEE Access*, vol. 13, pp. 92833-92847, 2025.
3. Sathvik Udupa, Jesuraja Bandekar, Abhayjeet Singh, Deekshitha G, Saurabh Kumar, Sandhya Badiger, Amala Nagireddi, Roopa R, Prasanta Kumar Ghosh, Hema A. Murthy, Pranaw Kumar, Keiichi Tokuda, Mark Hasegawa-Johnson, Philipp Olbrich, “LIMITS’24: Multi-Speaker, Multi-Lingual INDIC TTS With Voice Cloning” in *IEEE Open Journal of Signal Processing*, vol. 6, pp. 293-302, 2025.
4. Abhayjeet Singh, Amala Nagireddi, Anjali Jayakumar, Deekshitha G, Jesuraja Bandekar, Roopa R, Sandhya Badiger, Sathvik Udupa, Saurabh Kumar, Prasanta Kumar Ghosh, Hema A Murthy, Heiga Zen, Pranaw Kumar, Kamal Kant, Amol Bole, Bira Chandra Singh, Keiichi Tokuda, Mark Hasegawa-Johnson, and Philipp Olbrich, “Lightweight, Multi-Speaker, Multi-Lingual Indic Text-to-Speech,” in *IEEE Open Journal of Signal Processing*, vol. 5, pp. 790-798, 2024.
5. Yukiya Hono, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “PeriodNet: A Non-Autoregressive Raw Waveform Generative Model with a Structure Separating Periodic and Aperiodic Components,” *IEEE Access*, vol. 9, pp. 137599-137612, October, 2021.
6. Yukiya Hono, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Sinsy: A Deep Neural Network-Based Singing Voice Synthesis System,” *IEEE/ACM Transactions on Audio, Speech and Language Processing*, vol. 29, pp. 2803-2815, August, 2021.
7. Xin Wang, Shinji Takaki, Junichi Yamagishi, Simon King, and Keiichi Tokuda, “A vector quantized variational autoencoder (VQ-VAE) autoregressive neural F0 model for statistical parametric speech synthesis,” *IEEE/ACM Transactions on Audio, Speech, and Language Processing*, vol.28, pp.157–170, October 2019.
8. Kei Sawada, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “A Bayesian framework for image recognition based on hidden Markov eigen-image models,” *IEEE Transactions on Electrical and Electronic Engineering*, vol.13, Issue 9, pp.1335–1347, September 2018.
9. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Mel-cepstrum-based quantization noise shaping applied to neural-network-based speech waveform synthesis,” *IEEE/ACM Transactions on Audio, Speech and Language Processing*, vol. 26, Issue 7, pp.1173–1180, July 2018.

10. Kei Sawada, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Constructing text-to-speech systems for languages with unknown pronunciations," *Acoustical Science and Technology*, vol.39, Issue 2, pp.119–129, March 2018.
11. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Simultaneous optimization of multiple tree-based factor analyzed HMM for speech synthesis," *IEEE/ACM Transactions on Audio, Speech and Language Processing*, vol. 25, Issue 9, pp.1532–1541, September 2017.
12. Kei Sawada, Akira Tamamori, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "A Bayesian approach to image recognition based on separable lattice hidden Markov models," *IEICE Transactions on Information and Systems*, vol.E99-D, no.12, pp.3119–3131, December 2016.
13. Kazuhiro Nakamura, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, "Hidden Markov model-based English singing voice synthesis," *IEICE Transactions on Information and Systems*, vol.J97-D, no.10, pp.1572–1581, October 2014 (in Japanese).
14. Akira Tamamori, Yoshihiko Nankaku, Keiichi Tokuda, "Image recognition based on separable lattice trajectory 2-D HMMs," *IEICE Transactions on Information and Systems*, vol.E97-D, no.7, pp.1842–1854, July 2014.
15. Hongwu Yang, Keiichiro Oura, Haiyan Wang, Zhenye Gan, Keiichi Tokuda, "Using speaker adaptive training to realize Mandarin-Tibetan cross-lingual speech synthesis," *Springer, Multimedia Tools and Applications*, June 2014.
16. Kazuhiro Nakamura, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "Integration of spectral feature extraction and modeling for HMM-based speech synthesis," *IEICE Transactions on Information and Systems*, vol.E97-D, no.6, pp.1438–1448, June 2014.
17. Shinji Takaki, Yoshihiko Nankaku and Keiichi Tokuda, "Contextual additive structure for HMM-based speech synthesis," *IEEE Journal of Selected Topics in Signal Processing*, vol.8, issue 2, pp.229–238, April 2014.
18. Sayaka Shiota, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "A Bayesian framework using multiple model structures for speech recognition," *IEICE Transactions on Information and Systems*, vol.E96-D, no.4, pp.939–948, April 2013.
19. John Dines, Hui Liang, Lakshmi Saheer, Matthew Gibson, William Byrne, Keiichiro Oura, Keiichi Tokuda, Junichi Yamagishi, Simon King, Mirjam Wester, Teemu Hirsimäki, Reima Karhila, Mikko Kurimo, "Personalising speech-to-speech translation: unsupervised cross-lingual speaker adaptation for HMM-based speech synthesis," *Computer Speech and Language*, vol.27, pp.420–437, February 2013.
20. Kei Hashimoto, Junichi Yamagishi, William Byrne, Simon King, Keiichi Tokuda, "Impacts of machine translation and speech synthesis on speech-to-speech translation," *Speech Communication*, vol.54, no.7, pp.857–866, September 2012.
21. Akira Tamamori, Yoshihiko Nankaku, Keiichi Tokuda, "An extension of separable lattice 2-D HMMs for rotational data variations," *IEICE Transactions on Information and Systems*, vol.E95-D, no.8, pp.2074–2083, August 2012.

22. Keiichiro Oura, Junichi Yamagishi, Mirjam Wester, Simon King, Keiichi Tokuda, "Analysis of unsupervised cross-lingual speaker adaptation for HMM-based speech synthesis using KLD-based transform mapping," *Speech Communication*, vol.54, no.6, pp.703–714, July 2012.
23. Sayaka Shiota, Kei Hashimoto, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Speech recognition based on statistical models including multiple phonetic decision trees" *Acoustical Science and Technology*. (accepted)
24. Kei Hashimoto, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Bayesian context clustering using cross validation for speech recognition" *IEICE Transactions on Information and Systems*, vol.E94-D, no.3, pp.668–678, March 2011.
25. Ryuta Terashima, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "A frame-based context-dependent acoustic modeling for speech recognition," *The transactions of the Institute of Electrical Engineers of Japan C (A publication of Electronics, Information and System Society)*, vol.130, no.4, pp.557–564, April 2010 (in Japanese).
26. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Continuous stochastic feature mapping based on trajectory HMMs," *IEEE Transactions on Audio, Speech, and Language Processing*, vol.18, no.5, pp.417–430, Feb. 2011.
27. Ryuta Terashima, Takayoshi Yoshimura, Toshihiro Wakita, Keiichi Tokuda, Tadashi Kitamura, "Prediction method of speech recognition performance based on HMM-based speech synthesis technique," *The transactions of the Institute of Electrical Engineers of Japan C (A publication of Electronics, Information and System Society)*, vol.130, no.4, pp.557–564, April 2010 (in Japanese).
28. Keiichiro Oura, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "A covariance-tying technique for HMM-based speech synthesis," *IEICE Transactions on Information and Systems*, vol.E93-D, no.3, pp.595–601, March 2010.
29. Junichi Yamagishi, Bela Usabaev, Simon King, Oliver Watts, John Dines, Jilei Tian, Yong Guan, Rile Hu, Keiichiro Oura, Yi-Jian Wu, Keiichi Tokuda, Reima Karhila, Mikko Kurimo "Thousands of Voices for HMM-Based Speech Synthesis–Analysis and Application of TTS Systems Built on Various ASR Corpora," *IEEE Transactions on Audio, Speech, and Language Processing*, vol.18, no.5, pp.984–1004, July 2010.
30. Kei Hashimoto, Hirohumi Yamamoto, Hideo Okuma, Eiichiro Sumita, Keiichi Tokuda, "A reordering model using a source-side parse-tree for statistical machine translation," *IEICE Transactions on Information and Systems*, vol.E92-D, no.12, pp.2386–2393, December 2009.
31. Junichi Yamagishi, Takashi Nose, Heiga Zen, Zhen-Hua Ling, Tomoki Toda, Keiichi Tokuda, Simon King, Steve Renals, "A robust speaker-adaptive HMM-based text-to-speech synthesis," *IEEE Transactions on Audio, Speech, and Language Processing*, vol.17, no.6, August 2009.
32. Heiga Zen, Keiichi Tokuda, Alan W. Black, "Statistical parametric speech synthesis," *Speech Communication*, vol.51, no.11, pp.1039–1154, November 2009.

33. Keiichiro Oura, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "A Fully Consistent Hidden Semi-Markov Model-Based Speech Recognition System," IEICE Transactions on Information and Systems, vol.E91-D, no.11, pp.2693–2700, Nov 2008.
34. Heiga Zen, Tomoki Toda, Keiichi Tokuda, "The Nitech-NAIST HMM based speech synthesis system for the Blizzard Challenge 2006," IEICE Transactions on Information and Systems, vol.E91-D, no.6, pp.1764–1773, June 2008.
35. Tomoki Toda, Alan W. Black, and Keiichi Tokuda, "Statistical mapping between articulatory movements and acoustic spectrum with a Gaussian mixture model," Speech Communication, vol.50, no.3, pp.215–227, March 2008.
36. Tomoki Toda, Alan W. Black, and Keiichi Tokuda, "Voice conversion based on maximum likelihood estimation of speech parameter trajectory," IEEE Transactions on Audio, Speech and Language Processing, vol.15, no.8, pp.2222–2235, November 2007.
37. Ranniere Maia, Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, and Fernando G. V. Resende, "An HMM-based Brazilian Portuguese speech synthesizer and its characteristics," Journal of Communication and Information Systems, vol.21, no.2, pp.132–145, Aug. 2006.
38. Heiga Zen, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, and Tadashi Kitamura, "Hidden Semi-Markov Model Based Speech Synthesis System," IEICE Transactions on Information Systems, vol.E90-D, no.5, pp.825–834, May 2007.
39. Tomoki Toda and Keiichi Tokuda, "Speech Parameter Generation Algorithm Considering Global Variance for HMM-Based Speech Synthesis," IEICE Transactions on Information Systems, vol.E90-D, no.5, pp.816–824, May 2007.
40. Heiga Zen, Tomoki Toda, Masaru Nakamura, and Keiichi Tokuda, "Details of Nitech HMM-based speech synthesis system for the Blizzard Challenge 2005," IEICE Transactions on Information and Systems, vol.E90-D, no.1, pp.325–333, Jan. 2007.
41. Hisashi Kawai, Tomoki Toda, Junichi Yamagishi, Toshio Hirai, Jinfu Ni, Nobuyuki Nishizawa, Minoru Tsuzaki, and Keiichi Tokuda, "XIMERA: a concatenative speech synthesis system with large scale corpora," IEICE Transactions on Information and Systems, J89-D-II, no.12, pp.2688–2698, Dec. 2006 (in Japanese).
42. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "Reformulating the HMM as a trajectory model by imposing explicit relationships between static and dynamic feature vector sequences," Computer Speech and Language, vol.21, no.1, pp.153–173, Jan. 2007.
43. Amaro Lima, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, Tadashi Kitamura, and Fernando G. Resende, "Applying sparse KPCA for feature extraction in speech recognition," IEICE Transactions on Information Systems, vol.E88-D, no.3, pp.401–409, March 2005.
44. Hiroyuki Suzuki, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, and Tadashi Kitamura, "Continuous speech recognition based on general factor dependent acoustic models," IEICE Transactions on Information Systems, vol.E88-D, no.3, pp.410–417, March 2005.

45. Hiroyoshi Yamamoto, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, and Tadashi Kitamura, "Parameter Sharing in Mixture of Factor Analyzers for Speaker Identification," IEICE Transactions on Information Systems, vol.E88-D, no.3, pp.419–424, March 2005.
46. Yohei Itaya, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, and Tadashi Kitamura, "Deterministic annealing EM algorithm in acoustic modeling for speaker and speech recognition," IEICE Transactions on Information Systems, vol.E88-D, no.3, pp.425–431, March 2005.
47. Amaro Lima, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "On the use of kernel PCA for feature extraction in speech recognition," IEICE Transactions on Information Systems, vol.E87-D, no.12, pp.2802–2811, Dec. 2004.
48. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "Decision tree based simultaneous clustering of phonetic contexts, dimensions, and state positions for acoustic modeling," IEICE Trans. Inf. & Syst., vol.87-D-II, no.8, pp.1593–1602, Aug. 2004 (in Japanese).
49. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "Incorporation of mixed excitation model and postfilter into HMM-based text-to-speech synthesis," IEICE Trans. Inf. & Syst., vol.87-D-II, no.8, pp.1565–1571, Aug. 2004. (in Japanese). Translation: Systems and Computers in Japan, John Wiley & Sons, Inc., vol.36, no.12, pp.43–50, Nov. 15, 2005.
50. Shinji Sako, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "A singing voice synthesis system based on hidden Markov model," Transactions of Information Processing Society of Japan, vol.45, no.3, pp.719–727, Mar. 2004 (in Japanese).
51. Toru Takahashi, Keiichi Tokuda, Takao Kobayashi, Tadashi Kitamura, "Mixture density models based on mel-cepstral representation of Gaussian process," IEICE Trans. Fundamentals, vol.E86-A, no.8, pp.1971–1978, Aug. 2003.
52. Junichi Yamagishi, M.Tamura, T. Masuko, K.Tokuda and T.Kobayashi, "A training method of average voice model for HMM-based speech synthesis," IEICE Trans. Fundamentals of Electronics, Communications and Computer Sciences, E86-A, vol.8, pp.1956–1963 August 2003.
53. Junichi Yamagishi, Masatsune Tamura, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "A context clustering technique for average voice models," IEICE Trans. Inf. & Syst., vol.E86-D, no.3, pp.534–542, Mar. 2003.
54. Yoshihiko Nankaku, Keiichi Tokuda, Tadashi Kitamura, Takao Kobayashi, "Normalized training for HMM-based visual speech recognition," IEICE Trans. Inf. & Syst., vol.J86-D-II, no.2, pp.163–172, Feb. 2003 (in Japanese).
55. Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "Very Low Bit Rate Speech Coding Based on HMM with Speaker Adaptation," IEICE Trans. Inf. & Syst., vol.J85-D-II, no.12, pp.1749–1759, Dec. 2002 (in Japanese).
56. Takayuki Satoh, Takashi Masuko, Takao Kobayashi, Keiichi Tokuda, "Discrimination of synthetic speech generated by an HMM-based speech synthesis system for speaker verification," IPSJ Journal, vol.43, no.7, pp.2197–2204, July 2002 (in Japanese).

57. Shinichi Kawamoto, Hiroshi Shimodaira, Tsuneo Nitta, Takuya Nishimoto, Satoshi Nakamura, Kazunori Itou, Shigeo Morishima, Tatsuo Yotsukura, Akihiko Kai, Akinobu Lee, Yohichi Yamashita, Takao Kobayashi, Keiichi Tokuda, Keikichi Hirose, Nobuaki Mine-matsu, Atsushi Yamada, Yasuharu Den, Takenori Utsuro, Shigeki Sagayama, "Implementation of anthropomorphic spoken dialog agent," Transactions of Information Processing Society of Japan, vol.43, no.7, pp.2249-2263, July 2002 (in Japanese).
58. Sinji Sako, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "HMM-based audio-visual speech synthesis —Pixel-based approach," IPSJ Journal, vol.43, no.7, pp.2169-2176, July 2002 (in Japanese).
59. Mitsuru Katsumata, Hisashi Suzuki, Keiichi Tokuda, and Tadashi KITAMURA, "A Hierarchical Algorithm for Monotonic and Continuous Two-Dimensional Dynamic Programming," IEICE Trans. Inf. & Syst., vol.J85-D-II, no.9, pp.1382-1391, Sep. 2002 (in Japanese).
60. Toru Takahashi, Keiichi Tokuda, Takao Kobayashi, Tadashi Kitamura, "Vector quantization of mel-cepstrum coefficients using distortion measure for spectral analysis," IEICE Trans. Inf. & Syst., vol.J85-D-II, no.8, pp.1273-1283, Aug. 2002 (in Japanese).
61. Masatsune Tamura, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "Speaker adaptation of pitch and spectrum for HMM-based speech synthesis," IEICE Trans. Inf. & Syst., vol.J85-D-II, no.4, pp.545-553, Apr. 2002 (in Japanese).
62. Kazuhito Koishida, Keiichi Tokuda, Takashi Masuko and Takao Kobayashi, "Vector quantization of speech spectral parameters using statistics of static and dynamic features," IEICE Trans. Inf. & Syst., vol.E84-D, no.10, pp.1427-1434, Oct. 2001.
63. Chiyomi Miyajima, Y. Hattori, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Text-independent speaker identification using Gaussian Mixture models based on multi-space probability distribution," IEICE Trans. Inf. & Syst., vol.E84-D, no.7, pp.847-855, July 2001.
64. Chiyomi Miyajima, Hideyuki Watanabe, Keiichi Tokuda, Tadashi Kitamura, Shigeru Katagiri, "A New Approach to Designing a Feature Extractor in Speaker Identification Based on Discriminative Feature Extraction," Speech Communication, vol.35, no.3-4, pp.203-218, Oct. 2001.
65. Takashi Masuko, Masatsune Tamura, Keiichi Tokuda, and Takao Kobayashi, "Voice characteristics conversion for HMM-based speech synthesis system using MAP-VFS," IEICE Trans. Inf. & Syst., vol.J83-D-II, no.12, pp.2509-2516, Dec. 2000 (in Japanese).
66. Junibakti Sanubari and Keiichi Tokuda, "RLS-type two dimensional adaptive filter with a t-distribution assumption," Signal Processing, vol.80, no.12, pp.2483-2495, Nov. 2000.
67. Takashi Masuko, Keiichi Tokuda, and Takao Kobayashi, "Imposture against a speaker verification system using synthetic speech," IEICE Trans., vol.J83-D-II, no.11, pp.2283-2290, Nov. 2000 (in Japanese).
68. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Simultaneous Modeling of Spectrum, Pitch and Duration in HMM-based Speech Synthesis," IEICE Trans. Inf. & Syst., vol.J83-D-II, no.11, pp.2099-2107, Nov. 2000 (in Japanese).

69. Oscar Vanegas, Keiichi Tokuda, Tadashi Kitamura, "Lip location normalized training for visual speech recognition," *IEICE Trans. Inf. & Syst.*, vol.E83-D, no.11, pp.1969–1977, Nov. 2000.
70. Junibakti Sanubari and Keiichi Tokuda, "A new robust two dimensional spectral estimation based on an AR model excited by a t-distribution process and its QR-decomposition recursive algorithm," *Journal of Circuits, Systems, and Computers*, vol.9, nos.1–2, pp.51–66, Jan. 1999.
71. Takashi Masuko, Keiichi Tokuda, Noboru Miyazaki and Takao Kobayashi, "Pitch pattern generation using multi-space probability distribution HMM," *IEICE Trans.*, vol.J83-D-II, no.7, pp.1600–1609, July 2000 (in Japanese).
72. Keiichi Tokuda, Takashi Masuko, Noboru Miyazaki and Takao Kobayashi, "Multi-space Probability distribution HMM," *IEICE Trans.*, vol.J83-D-II, no.7, pp.1579–1589, July 2000 (in Japanese).
73. Kazuhito Koishida, Goh Hirabayashi, Keiichi Tokuda, and Takao Kobayashi, "A 16kbit/s wideband CELP-based speech coder using mel-generalized cepstral analysis," *IEICE Trans. Inf. & Syst.*, vol.E83-D, no.4, pp.876–883, Apr. 2000.
74. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Speaker interpolation for HMM-based speech synthesis system," *The Journal of the Acoustical Society of Japan (E)*, vol.21, no.4, pp.199–206, Apr. 2000.
75. Takeshi Wakako, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Speech spectral estimation based on expansion of log spectrum by arbitrary basis functions," *IEICE Trans.*, vol.J82-D-II, no.12, pp.2203–2211, Dec. 1999.
76. Jun Hiroi, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Very low bit rate speech coding based on HMMs," *IEICE Trans.*, vol.J82-D-II, no.11, pp.1857–1864, Nov. 1999.
77. Fernando G. Resende, Paulo S. R. Diniz, Keiichi Tokuda, Mineo Kaneko and Akinori Nishihara, "New adaptive algorithms based on multi-band decomposition of the error signal," *IEEE Trans. Circuits and Systems II*, vol.45, no.5, pp.592–599, May 1998.
78. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "CELP speech coding based on mel-generalized cepstral analysis," *IEICE Trans.*, vol.J81-A, no.2, pp.252–260, Feb. 1998 (in Japanese).
79. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Spectral representation of speech based on mel-generalized cepstral coefficients and its properties," *IEICE Trans.*, vol.J80-A, no.11, pp.1999–2006, Nov. 1997 (in Japanese).
80. Fernando G. Resende, Paulo S. R. Diniz, Keiichi Tokuda, Mineo Kaneko and Akinori Nishihara, "LMS-based algorithms with multi-band decomposition of the estimation error applied to system identification," *IEICE Trans.*, vol.E80-A, no.8, pp.1376–1383, Aug. 1997.
81. Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Satoshi Imai, "An algorithm for speech parameter generation from HMM using dynamic features," *the Journal of the Acoustical Society of Japan*, vol.53, no.3, pp.192–200, Mar. 1997 (in Japanese).

82. Fernando G. Resende, Keiichi Tokuda, Mineo Kaneko and Akinori Nishihara, "Multi-band decomposition of the linear prediction error applied to adaptive AR spectral estimation," *IEICE Trans.*, vol.E80-A, no.2, pp.365–376, Feb. 1997.
83. Takashi Masuko, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "HMM-based speech synthesis using dynamic features," *IEICE Trans.*, vol.J79-D-II, no.12, pp.2184–2190, Dec. 1996 (in Japanese).
84. Fernando G. Resende, Keiichi Tokuda and Mineo Kaneko, "Adaptive AR spectrum estimation based on wavelet decomposition of linear prediction error," *IEICE Trans.*, vol.E79-A, no.5, pp.665–673, May 1996.
85. Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Adaptive cepstral analysis of speech," *IEEE Trans. Speech and Audio Proces.*, vol.SA-3, no.6, pp.481–489, Nov. 1995.
86. Yukihiro Naito, Keiichi Tokuda and Mineo Kaneko, "An adaptive algorithm based on non-uniform resolution filter," *IEICE Trans.*, vol.J78-A, no.9, pp.1092–1102, Sep. 1995 (in Japanese). Translation: *Electronics and Communications in Japan, Part 3, Scripta Technica, Inc.*, vol.79, no.9, pp.54–66, 1996.
87. Keiichi Tokuda, Takao Kobayashi, Toshiaki Fukada and Satoshi Imai, "Speech coding based on adaptive mel-cepstral analysis and its evaluation," *IEICE Trans.*, vol.J77-A, no.11, pp.1443–1452, Nov. 1994 (in Japanese). Translation: *Electronics and Communications in Japan, Part 3, Scripta Technica, Inc.*, vol.78, no.6, pp.50–61, 1995.
88. Junibakti Sanubari, Keiichi Tokuda, Mahoki Onoda, "Time series analysis based on exponential model excited by t-distribution process and its algorithm," *IEICE Trans.*, vol.E76-A, no.5, pp.808–819, May 1993.
89. Junibakti Sanubari, Keiichi Tokuda, Mahoki Onoda, "Speech analysis based on AR model driven by t-distribution process," *IEICE Trans.*, vol.E75-A, no.9, pp.1159–1169, Sep. 1992.
90. Keiichi Tokuda, Takao Kobayashi, Kenji Chiba and Satoshi Imai, "Spectral estimation of speech by mel-generalized cepstral analysis," *IEICE Trans.*, vol.J75-A, no.7, pp.1124–1134, July 1992 (in Japanese). Translation: *Electronics and Communications in Japan, Part 3, Scripta Technica, Inc.*, vol.76, no.2, pp.30–43, 1993.
91. Takao Kobayashi, Kazuyoshi Fukushi, Keiichi Tokuda and Satoshi Imai, "2-D LMA filters —design of stable two-dimensional digital filters," *IEICE Trans.*, vol.E75-A, no.2, pp.240–246, Feb. 1992.
92. Keiichi Tokuda, Takao Kobayashi, Toshiaki Fukada and Satoshi Imai, "Adaptive mel-cepstral analysis of speech," *IEICE Trans.*, vol.J74-A, no.8, pp.1249–1256, Aug. 1991 (in Japanese).
93. Keiichi Tokuda, Takao Kobayashi, Toshiaki Fukada, Hironori Saito and Satoshi Imai, "Spectral estimation of speech based on Mel-cepstral representation," *IEICE Trans.*, vol.J74-A, no.8, pp.1240–1248, Aug. 1991 (in Japanese).

94. Keiichi Tokuda, Takao Kobayashi, Soji Shiimoto and Satoshi Imai, “Adaptive cepstral analysis —Adaptive filtering based on cepstral representation—,” IEICE Trans., vol.J73-A, no.7, pp.1207–1215, July 1990 (in Japanese). Translation: Electronics and Communications in Japan, Part 3, Scripta Technica, Inc., vol.74, no.3, pp.36–45, 1991.
95. Keiichi Tokuda, Takao Kobayashi, Ryutaro Yamamoto and Satoshi Imai, “Spectral estimation of speech based on generalized cepstral representation,” IEICE Trans., vol.J72-A, no.3, pp.457–465, Mar. 1989 (in Japanese). Translation: Electronics and Communications in Japan, Part 3, Scripta Technica, Inc., vol.73, no.1, pp.72–81, 1990.
96. Keiichi Tokuda, Takao Kobayashi, Atsuhiko Tokuda and Satoshi Imai, “Simultaneous approximation of magnitude and phase for digital filters by decomposition of maximum and minimum phase components,” IEICE Trans., vol.J71-A, no.2, pp.260–267, Feb. 1988 (in Japanese).
97. Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, “Cepstral analysis with non-uniform spectral weighting for spectral envelope extraction,” IEICE Trans., vol.J70-A, no.6, pp.952–959, June 1987 (in Japanese). Translation: Electronics and Communications in Japan, Part 3, Scripta Technica, Inc., vol.72, no.1, pp.20–28, 1989.

REFEREED JOURNAL LETTERS

1. Heiga Zen, Takashi Masuko, Keiichi Tokuda, Takayoshi Yoshimura, Takao Kobayashi, and Tadashi Kitamura, “State Duration Modeling for HMM-Based Speech Synthesis,” IEICE Transactions on Information Systems, vol.E90-D, no.3, pp.692–693, Mar. 2007.
2. Junibakti Sanubari, Keiichi Tokuda, Mahoki Onoda, “Robust recursive time series modeling based on an AR model excited by a t-distribution process,” IEEE Trans. Signal Proces., vol.46, no.1, pp.218–222, Jan. 1998.
3. Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, “Recursion formula for calculation of mel generalized cepstrum coefficients,” IEICE Trans., vol.J71-A, no.1, pp.128–131, Jan. 1988 (in Japanese).

REFEREED CONFERENCE AND WORKSHOP PAPERS

1. Takenori Yoshimura, Shinji Takaki, Kazuhiro Nakamura, Keiichiro Oura, Takato Fujimoto, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, “SSLZip: Simple autoencoding for enhancing self-supervised speech representations in speech generation,” 13th ISCA Speech Synthesis Workshop, pp. 117–122, Leeuwarden, Netherlands, August, 2025.
2. Masato Takagi, Miku Nishihara, Yukiya Hono, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, “PeriodCodec: A Pitch-Controllable Neural Audio Codec Using Periodic Signals for Singing Voice Synthesis,” Interspeech 2025, pp. 4913–4917, Rotterdam, Netherlands, August 17–21, 2025.
3. Miku Nishihara, Dan Wells, Korin Richmond and Aidan Pine, “Low-dimensional Style Token Control for Hyperarticulated Speech Synthesis,” Interspeech 2024, pp. 3385–3389, Greece, September, 2024.

4. Yukiya Hono, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "PeriodGrad: Towards Pitch-Controllable Neural Vocoder Based on a Diffusion Probabilistic Model," 2024 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), pp. 12782-12786, Korea, April, 2024.
5. Abhayjeet Singh, Amala Nagireddi, Deekshitha G, Jesuraja Bandekar, Roopa R, Sandhya Badiger, Sathvik Udupa, Prasanta Kumar Ghosh, Hema A Murthy, Pranaw Kumar, Keiichi Tokuda, Mark Hasegawa-Johnson, and Philipp Olbrich, "LIMMITS ' 24: Multi-Speaker, Multi-Lingual Indic TTS with Voice Cloning," 2024 IEEE International Conference on Acoustics, Speech, and Signal Processing Workshops (ICASSPW), pp. 61-62, Seoul, Korea, Republic of, 2024.
6. Takenori Yoshimura, Takato Fujimoto, Keiichiro Oura, and Keiichi Tokuda, "SPTK4: An open-source software toolkit for speech signal processing," 12th ISCA Speech Synthesis Workshop, pp. 211-217, Grenoble, France, August, 2023.
7. Yukiya Hono, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Singing voice synthesis based on a musical note position-aware attention mechanism," 2023 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), Greece, June, 2023.
8. Takenori Yoshimura, Shinji Takaki, Kazuhiro Nakamura, Keiichiro Oura, Yukiya Hono, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Embedding a differentiable mel-cepstral synthesis filter to a neural speech synthesis system," 2023 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), Greece, June, 2023.
9. Abhayjeet Singh, Amala Nagireddi, Deekshitha G, Jesuraja Bandekar, Roopa R, Sandhya Badiger, Sathvik Udupa, Prasanta Kumar Ghosh, Hema A Murthy, Heiga Zen, Pranaw Kumar, Kamal Kant, Amol Bole, Bira Chandra Singh, Keiichi Tokuda, Mark Hasegawa-Johnson, and Philipp Olbrich, "Lightweight, Multi-Speaker, Multi-Lingual Indic Text-to-Speech," ICASSP 2023 - 2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), pp. 1-2, 2023.
10. Kentaro Mitsui, Tianyu Zhao, Kei Sawada, Yukiya Hono, Yoshihiko Nankaku, and Keiichi Tokuda, "End-to-End Text-to-Speech Based on Latent Representation of Speaking Styles Using Spontaneous Dialogue," Interspeech 2022, pp. 2328-2332, Incheon, Korea, September, 2022.
11. Takato Fujimoto, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Autoregressive variational autoencoder with a hidden semi-Markov model-based structured attention for speech synthesis," 2022 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), pp. 7462-7466, Singapore, Singapore, May, 2022.
12. Yukiya Hono, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "PeriodNet: A Non-Autoregressive Raw Waveform Generative Model with a Structure Separating Periodic and Aperiodic Components," IEEE Access, vol. 9, pp. 137599-137612, October, 2021.
13. Yukiya Hono, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Sinsy: A Deep Neural Network-Based Singing Voice Synthesis System," IEEE/ACM Transactions on Audio, Speech and Language Processing, vol. 29, pp. 2803-2815, August, 2021.

14. Yukiya Hono, Kazuna Tsuboi, Kei Sawada, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Hierarchical Multi-Grained Generative Model for Expressive Speech Synthesis," *Interspeech 2020*, pp. 3441–3445, Shanghai, China, October, 2020.
15. Takato Fujimoto, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Semi-supervised learning based on hierarchical generative models for end-to-end speech synthesis," *2020 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP)*, pp. 7644–7648, Barcelona, Spain, May, 2020.
16. Kazuhiro Nakamura, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Fast and high-quality singing voice synthesis system based on convolutional neural networks," *2020 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP)*, pp. 7239–7243, Barcelona, Spain, May, 2020.
17. Motoki Shimada, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Low computational cost speech synthesis based on deep neural networks using hidden semi-Markov model structures," *10th ISCA Speech Synthesis Workshop (SSW10)*, pp.177–182, Vienne, Austria, September, 2019.
18. Takato Fujimoto, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Impacts of input linguistic feature representation on Japanese end-to-end speech synthesis," *10th ISCA Speech Synthesis Workshop (SSW10)*, pp.166–171, Vienne, Austria, September, 2019.
19. Keiichiro Oura, Kazuhiro Nakamura, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Deep neural network based real-time speech vocoder with periodic and aperiodic inputs," *10th ISCA Speech Synthesis Workshop (SSW10)*, pp.13–18, Vienne, Austria, September, 2019.
20. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Speaker-dependent WaveNet-based delay-free ADPCM speech coding," *2019 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2019)*, pp.7145–7149, Brighton, UK, May, 2019.
21. Yukiya Hono, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Singing voice synthesis based on generative adversarial networks," *2019 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2019)*, pp.6955–6959, Brighton, UK, May, 2019.
22. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "WaveNet-based zero-delay lossless speech coding," *2018 IEEE Workshop on Spoken Language Technology (SLT 2018)*, pp.153–158, Athens, Greece, December 18–21, 2018.
23. Kento Nakao, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Speaker adaptation for speech synthesis based on deep neural networks using hidden semi-Markov model structures," *Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018)*, pp.638–643, Honolulu, Hawaii, November 12–15, 2018.

24. Takayuki Kasugai, Yoshinari Tsuzuki, Kei Sawada, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Image recognition based on convolutional neural networks using features generated from separable lattice hidden Markov models,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018), pp.324–328, Honolulu, Hawaii, November 12–15, 2018.
25. Koki Senda, Yukiya Hono, Kei Sawada, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Singing voice conversion using posted waveform data on music social media,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018), pp.1913–1917, Honolulu, Hawaii, November 12–15, 2018.
26. Yukiya Hono, Shumma Murata, Kazuhiro Nakamura, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Recent development of the DNN-based singing voice synthesis system – Sinsy,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018), pp.1003–1009, Honolulu, Hawaii, November 12–15, 2018.
27. Takenori Yoshimura, Natsumi Koike, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Discriminative feature extraction based on sequential variational autoencoder for speaker recognition,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018), pp.1742–1746, Honolulu, Hawaii, November 12–15, 2018.
28. Takato Fujimoto, Takenori Yoshimura, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Speech synthesis using WaveNet vocoder based on periodic/aperiodic decomposition,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2018 (APSIPA ASC 2018), pp.644–648, Honolulu, Hawaii, November 12–15, 2018.
29. Kei Sawada, Takenori Yoshimura, Kei Hashimoto, Keiichi Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “The NITech text-to-speech system for the Blizzard Challenge 2018,” Blizzard Challenge 2018 Workshop, Hyderabad, India, September 8, 2018. (web proceedings)
30. Eiji Ichikawa, Kei Sawada, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, “Image recognition based on separable lattice HMMs using a deep neural network for output probability distributions,” 2018 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2018), pp.3021–3025, Calgary, Canada, April 15–20, 2018.
31. Jumpei Niwa, Takenori Yoshimura, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, “Statistical voice conversion based on WaveNet,” 2018 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2018), pp.5289–5293, Calgary, Canada, April 15–20, 2018.
32. Kei Sawada, Keiichi Tokuda, Simon King, and Alan W Black, “The Blizzard Machine Learning Challenge 2017,” 2017 IEEE Automatic Speech Recognition and Understanding Workshop (ASRU 2017), pp.331–337, Okinawa, Japan, December, 2017.

33. Kei Sawada, Kei Hashimoto, Keiichiro Oura, and Keiichi Tokuda, "The NITech text-to-speech system for the Blizzard Challenge 2017," *Blizzard Challenge 2017 Workshop*, Stockholm, Sweden, August 25, 2017. (web proceedings)
34. Amelia J. Gully, Takenori Yoshimura, Damian Murphy, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Articulatory text-to-speech synthesis using the digital waveguide mesh driven by a deep neural network," *Interspeech 2017*, pp.234–238, Stockholm, Sweden, August 20–24, 2017.
35. Yoshinari Tsuzuki, Kei Sawada, Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Image recognition based on discriminative models using features generated from separable lattice HMMs," *2017 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2017)*, pp.2607–2611, New Orleans, USA, March 5–9, 2017.
36. Kei Sawada, Chiaki Asai, Kei Hashimoto, Keiichiro Oura, and Keiichi Tokuda, "The NITech text-to-speech system for the Blizzard Challenge 2016," *Blizzard Challenge 2016 Workshop*, California, USA, September 16, 2016. (web proceedings)
37. Keiichi Tokuda, Kei Hashimoto, Keiichiro Oura, and Yoshihiko Nankaku, "Temporal modeling in neural network based statistical parametric speech synthesis," *9th ISCA Speech Synthesis Workshop (SSW9)*, pp.113–118, California, USA, September 13–15, 2016.
38. Rasmus Dall, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Redefining the linguistic context feature set for HMM and DNN TTS through position and parsing," *Interspeech 2016*, pp.2851–2855, California, USA, September 8–12, 2016.
39. Takenori Yoshimura, Gustav Eje Henter, Oliver Watts, Mirjam Wester, Junichi Yamagishi, and Keiichi Tokuda, "A hierarchical predictor of synthetic speech naturalness using neural networks," *Interspeech 2016*, pp.342–346, California, USA, September 8–12, 2016.
40. Masanari Nishimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Singing voice synthesis based on deep neural networks," *Interspeech 2016*, pp.2478–2482, California, USA, September 8–12, 2016.
41. Naoki Hosaka, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Voice conversion based on trajectory model training of neural networks considering global variance," *Interspeech 2016*, pp.307–311, California, USA, September 8–12, 2016.
42. Keiichi Tokuda, and Heiga Zen, "Directly modeling voiced and unvoiced components in speech waveforms by neural networks," *2016 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2016)*, pp.5640–5644, Shanghai, China, March, 2016.
43. Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "Trajectory training considering global variance for speech synthesis based on neural networks," *2016 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2016)*, pp.5600–5604, Shanghai, China, March 20–25, 2016.
44. Kei Sawada, Kei Hashimoto, Keiichiro Oura, and Keiichi Tokuda, "The NITECH HMM-based text-to-speech system for the Blizzard Challenge 2015," *Blizzard Challenge 2015 Workshop*, Berlin, Germany, September 11, 2015. (web proceedings)

45. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “Simultaneous optimization of multiple tree structures for factor analyzed HMM-based speech synthesis,” *Interspeech 2015*, pp.1196–1200, Dresden, Germany, September 6–10, 2015.
46. Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, “The effect of neural networks in statistical parametric speech synthesis,” *2015 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2015)*, pp.4455–4459, Brisbane, Australia, April 19–24, 2015.
47. Keiichi Tokuda and Heiga Zen, “Directly modeling speech waveforms by neural networks for statistical parametric speech synthesis,” *2015 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2015)*, pp.4215–4219, Brisbane, Australia, April, 2015.
48. Daisuke Yamamoto, Keiichiro Oura, Ryota Nishimura, Takahiro Uchirya, Akinobu Lee, Ichi Takumi, Keiichi Tokuda, “Voice interaction system with 3D-CG virtual agent for stand-alone smartphones,” *the 2nd International Conference on Human Agent Interaction (HAI 2014)*, ACM digital library, pp.320–330, October 2014.
49. Kei Sawada, Shinji Takaki, Kei Hashimoto, Keiichiro Oura, Keiichi Tokuda, “Overview of NITECH HMM-based text-to-speech synthesis system for Blizzard Challenge 2014,” *Blizzard Challenge 2014 Workshop*, Singapore, September 19, 2014 (web proceedings).
50. Kazuhiro Nakamura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, “A mel-cepstral analysis technique restoring high frequency components from low-sampling-rate speech,” *Interspeech 2014*, pp.2494–2498, Singapore, September 14–18, 2014.
51. Kanako Shirota, Kazuhiro Nakamura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, “Integration of speaker and pitch adaptive training for HMM-based singing voice synthesis,” *2014 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2014)*, pp.2578–2582, Florence, Italy, May 4–9, 2014.
52. Kazuhiro Nakamura, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, “HMM-based singing voice synthesis and its application to Japanese and English,” *2014 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2014)*, pp.265–269, Florence, Italy, May 4–9, 2014.
53. Kei Sawada, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, “Image recognition based on hidden Markov eigen-image models using variational Bayesian method,” *Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2013 (APSIPA ASC 2013)*, Kaohsiung, Taiwan, October 29–November 1, 2013.
54. Hong-wu Yang, Keiichiro Oura, Zhen-ye Gan, Keiichi Tokuda, “Realizing Tibetan speech synthesis by speaker adaptive training,” *Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2013 (APSIPA ASC 2013)*, Kaohsiung, Taiwan, October 29–November 1, 2013.

55. Shinji Takaki, Kei Sawada, Kei Hashimoto, Keiichiro Oura, Keiichi Tokuda, "Overview of NIT HMM-based speech synthesis system for Blizzard Challenge 2013," *Blizzard Challenge 2013 Workshop*, Barcelona, Spain, September 3, 2013 (web proceedings).
56. Takenori Yoshimura, Kei Hashimoto, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, "Cross-lingual speaker adaptation based on factor analysis using bilingual speech data for HMM-based speech synthesis," *8th ISCA Speech Synthesis Workshop*, pp.317–322, Lyon, France, August 31–September 2, 2013.
57. Takaya Makino, Shinji Takaki, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "Separable lattice 2-d HMMs introducing state duration control for recognition of images with various variations," *2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pp.3203–3207, Vancouver, Canada, May 26–31, 2013.
58. Akira Tamamori, Yoshihiko Nankaku, Keiichi Tokuda, "Image recognition based on separable lattice trajectory HMMs," *2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pp.3467–3471, Vancouver, Canada, May 26–31, 2013.
59. Shinji Takaki, Yoshihiko Nankaku, Keiichi Tokuda, "Contextual Partial Additive Structure for HMM-based Speech Synthesis," *2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pp.7878–7882, Vancouver, Canada, May 26–31, 2013.
60. Kazuhiro Nakamura, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "Integration of Acoustic Modeling and Mel-cepstral analysis for HMM-based Speech Synthesis," *2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pp.7883–7887, Vancouver, Canada, May 26–31, 2013.
61. Akinobu Lee, Keiichiro Oura, Keiichi Tokuda, "MMDAgent –A FULLY OPEN-SOURCE TOOLKIT FOR VOICE INTERACTION SYSTEMS–," *2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2013)*, pp.8382–8385, Vancouver, Canada, May 26–31, 2013.
62. Shinji Takaki, Kei Sawada, Kei Hashimoto, Keiichiro Oura, Keiichi Tokuda, "Overview of NIT HMM-based speech synthesis system for Blizzard Challenge 2012," *Blizzard Challenge 2012 Workshop*, Portland, USA, September 14, 2012 (web proceedings).
63. Takafumi Hattori, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "A Bayesian approach to speaker recognition based on GMMs using multiple model structures," *Interspeech 2012*, Portland, USA, September 9–13, 2012.
64. Viviane de Franca Oliveira, Sayaka Shiota, Yoshihiko Nankaku, Keiichi Tokuda, "Cross-lingual speaker adaptation for HMM-based speech synthesis based on perceptual characteristics and speaker interpolation," *Interspeech 2012*, Portland, USA, September 9–13, 2012.
65. Kei Sawada, Akira Tamamori, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "Face recognition based on separable lattice 2-D HMMs using variational Bayesian method," *2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2012)*, pp.2205–2208, Kyoto, Japan, March 25–30, 2012.

66. Keisuke Kumaki, Yoshihiko Nankaku, Keiichi Tokuda, "Face recognition based on separable lattice 2-D HMMs using variational Bayesian method," 2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2012), pp.2209–2212, Kyoto, Japan, March 25–30, 2012.
67. Sayaka Shiota, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "A model structure integration based on a Bayesian framework for speech recognition," 2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2012), pp.4813–4816, Kyoto, Japan, March 25–30, 2012.
68. Keiichiro Oura, Ayami Mase, Yoshihiko Nankaku, Keiichi Tokuda, "Pitch adaptive training for HMM-based singing voice synthesis," 2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2012), pp.5377–5380, Kyoto, Japan, March 25–30, 2012.
69. Kei Hashimoto, Shinji Takaki, Keiichiro Oura, Keiichi Tokuda, "Overview of NIT HMM-based speech synthesis system for Blizzard Challenge 2011," Blizzard Challenge 2011 Workshop, Turin, Italy, September 2, 2011 (web proceedings).
70. Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Multi-speaker modeling with shared prior distributions and model structures for Bayesian speech synthesis," Interspeech 2011, pp.113–116, Florence, Italy, August 28–31, 2011.
71. Minoru Tsuzaki, Keiichi Tokuda, Hisashi Kawai, Jinfu Ni, "Estimation of perceptual spaces for speaker identities based on the cross-lingual discrimination task," Interspeech 2011, pp.157–160, Florence, Italy, August 28–31, 2011.
72. Lei Li, Yoshihiko Nankaku, Keiichi Tokuda, "A Bayesian approach to voice conversion based on GMMs using multiple model structures," Interspeech 2011, pp.661–664, Florence, Italy, August 28–31, 2011.
73. Ulpu Remes, Yoshihiko Nankaku, Keiichi Tokuda, "GMM-based missing-feature reconstruction on multi-frame windows," Interspeech 2011, pp.1665–1668, Florence, Italy, August 28–31, 2011.
74. Ling-Hui Chen, Yoshihiko Nankaku, Heiga Zen, Keiichi Tokuda, Zhen-Hua Ling, Li-Rong Dai, "Estimation of window coefficients for dynamic feature extraction for HMM-based speech synthesis," Interspeech 2011, pp.1801–1804, Florence, Italy, August 28–31, 2011.
75. Tsuneo Kato, Makoto Yamada, Nobuyuki Nishizawa, Keiichiro Oura, Keiichi Tokuda, "Large-scale subjective evaluations of speech rate control methods for HMM-based speech synthesizers," Interspeech 2011, pp.1845–1848, Florence, Italy, August 28–31, 2011.
76. Naoaki Ito, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Evaluation of tree-trellis based decoding in over-million LVCSR," Interspeech 2011, pp.1937–1940, Florence, Italy, August 28–31, 2011.
77. Kei Hashimoto, Junichi Yamagishi, William Byrne, Simon King, and Keiichi Tokuda, "An analysis of machine translation and speech synthesis in speech-to-speech translation system," 2011 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2011), pp.5108–5111, Prague, Czech Republic, May 22–27, 2011.

78. Shinji Takaki, Keiichiro Oura, Yoshihiko Nankaku, and Keiichi Tokuda, "An Optimization Algorithm of Independent Mean and Variance Parameter Tying Structures for HMM-based speech synthesis," 2011 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2011), pp.5108–5111, Prague, Czech Republic, May 22–27, 2011.
79. Shifeng Pan, Yoshihiko Nankaku, Keiichi Tokuda, Jianhua Tao, "Global variance modeling on frequency domain delta LSP for HMM-based speech synthesis," 2011 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2011), pp.4716–4719, Prague, Czech Republic, May 22–27, 2011.
80. Xianglin Peng, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, "Cross-Lingual Speaker Adaptation for HMM-Based Speech Synthesis Considering Differences Between Language-Dependent Average Voices," Proc. of IEEE 10th International Conference on Signal Processing, pp.605–608, Beijing China, 24–28 Oct. 2010.
81. Toyohiro Hayashi, Yoshihiko Nankaku, Akinobu Lee and Keiichi Tokuda, "Speaker Adaptation Based on Nonlinear Spectral Transform for Speech Recognition," Interspeech 2009, pp.542-545, Chiba Japan, 26–30 Sep. 2010.
82. Ayami Mase, Keiichiro Oura, Yoshihiko Nankaku, Keiichi Tokuda, "HMM-based singing voice synthesis system using pitch-shifted pseudo training data," Interspeech 2009, pp.845-848, Chiba Japan, 26–30 Sep. 2010.
83. Akira Saito, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Voice activity detection based on conditional random fields using multiple features," Interspeech 2009, pp.2086-2089, Chiba Japan, 26–30 Sep. 2010.
84. Keiichiro Oura, Kei Hashimoto, Sayaka Shiota, and Keiichi Tokuda, "Overview of NIT HMM-based speech synthesis system for Blizzard Challenge 2010," Blizzard Challenge 2010 Workshop, Kyoto, Japan, September 25, 2010.
85. Shinji Takaki, Yoshihiko Nankaku, and Keiichi Tokuda, "Spectral modeling with contextual additive structure for HMM-based speech synthesis," Proc. of 7th ISCA Speech Synthesis Workshop, pp.100–105, Kyoto, Japan, Sep. 22–24, 2010.
86. Keiichiro Oura, Ayami Mase, Tomohiko Yamada, Satoru Muto, Yoshihiko Nankaku, and Keiichi Tokuda, "Recent Development of the HMM-based Singing Voice Synthesis System – Sinsy," Proc. of 7th ISCA Speech Synthesis Workshop, pp.211–216, Kyoto, Japan, Sep. 22–24, 2010.
87. Kei Hashimoto, Yoshihiko Nankaku, and Keiichi Tokuda, "Bayesian speech synthesis framework integrating training and synthesis processes," Proc. of 7th ISCA Speech Synthesis Workshop, pp.106–111, Kyoto, Japan, Sep. 22–24, 2010.
88. Mirjam Wester, John Dines, Matthew Gibson, Hui Liang, Yi-Jian Wu, Lakshmi Saheer, Simon King, Keiichiro Oura, Philip N. Garner, William Byrne, Yong Guan, Teemu Hirsimäki, Reima Karhila, Mikko Kurimo, Matt Shannon, Sayaka Shiota, Jilei Tian, Keiichi Tokuda, Junichi Yamagishi, "Speaker adaptation and the evaluation of speaker similarity in the EMIME speech-to-speech translation project," Proc. of 7th ISCA Speech Synthesis Workshop, pp.192–197, Kyoto, Japan, Sep. 22–24, 2010.

89. Yoshiaki Takahashi, Akira Tamamori, Yoshihiko Nankaku, Keiichi Tokuda, "Face recognition based on separable lattice 2-D HMM with state duration modeling," 2010 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2010), pp.2162–2165, Dallas, Texas, U.S.A., March 14–19, 2010.
90. Kyosuke Kazumi, Yoshihiko Nankaku, Keiichi Tokuda, "Factor analyzed voice models for HMM-based speech synthesis," 2010 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2010), pp.4234–4237, Dallas, Texas, U.S.A., March 14–19, 2010.
91. Akira Tamamori, Yoshihiko Nankaku, Keiichi Tokuda, "An extension of separable lattice 2-D HMMs for rotational data variations," 2010 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2010), pp.2206–2209, Dallas, Texas, U.S.A., March 14–19, 2010.
92. Heiga Zen, Mark Gales, Yoshihiko Nankaku, Keiichi Tokuda, "Statistical parametric speech synthesis based on product of experts," 2010 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2010), pp.4242–4245, Dallas, Texas, U.S.A., March 14–19, 2010.
93. Keiichiro Oura, Keiichi Tokuda, Junichi Yamagishi, Simon King, Mirjam Wester, "Unsupervised cross-lingual speaker adaptation for HMM-based speech synthesis," 2010 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2010), pp.4594–4597, Dallas, Texas, U.S.A., March 14–19, 2010.
94. Junichi Yamagishi, Bela Usabaev, Simon King, Oliver Watts, John Dines, Jilei Tian, Rile Hu, Yong Guan, Keiichiro Oura, Keiichi Tokuda, Reima Karhila, Mikko Kurimo "Thousands of Voices for HMM-based Speech Synthesis," Interspeech 2009, pp.420–423, Brighton, U.K., September 6-10, 2009.
95. Yi-Jian Wu, Yoshihiko Nankaku, Keiichi Tokuda, "State mapping based method for cross-lingual speaker adaptation in HMM-based speech synthesis," Interspeech 2009, pp.528–531, Brighton, U.K., September 6–10, 2009.
96. Sayaka Shiota, Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "Deterministic Annealing Based Training Algorithm for Bayesian Speech Recognition," Interspeech 2009, pp.680–683, Brighton, U.K., September 6-10, 2009.
97. Kei Hashimoto, Yoshihiko Nankaku, Keiichi Tokuda, "A Bayesian Approach to Hidden Semi-Markov Model Based Speech Synthesis," Interspeech 2009, pp.1751–1754, Brighton, U.K., September 6–10, 2009.
98. Keiichiro Oura, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Tying covariance matrices to reduce the footprint of HMM-based speech synthesis systems," Interspeech 2009, pp.1751–1762, Brighton, U.K., September 6–10, 2009.
99. Ranniery Maia, Tomoki Toda, Keiichi Tokuda, Shinsuke Sakai, Satoshi Nakamura, "A decision tree-based clustering approach to state definition in an excitation modeling framework for HMM-based speech synthesis," Interspeech 2009, pp.1783–1786, Brighton, U.K., September 6–10, 2009.

100. Yi-Jian Wu, Long Qin, Keiichi Tokuda, "An improved minimum generation error based model adaptation for HMM-based speech synthesis," *Interspeech 2009*, pp.1787–1790, Brighton, U.K., September 6–10, 2009.
101. Heiga Zen, Keiichiro Oura, Takashi Nose, Junichi Yamagishi, Shinji Sako, Tomoki Toda, Takashi Masuko, Alan W. Black, Keiichi Tokuda, "Recent development of the HMM-based speech synthesis system (HTS)," *Asia-Pacific Signal and Information Processing Association 2009 Annual Summit and Conference (APSIPA ASC 2009)*, pp.121–130, Sapporo, Japan, October 4–7, 2009.
102. Keiichiro Oura, Yi-Jian Wu, Keiichi Tokuda, "Overview of NIT HMM-based speech synthesis system for Blizzard Challenge 2009," *2009 Blizzard Challenge Workshop*, September 4, 2009 (web proceedings).
103. Kei Hashimoto, Hirohumi Yamamoto, Hideo Okuma, Eiichiro Sumita, Keiichi Tokuda, "Reordering model using syntactic information of a source tree for statistical machine translation," *NAACL HLT 2009 Workshop: Third Workshop on Syntax and Structure in Statistical Translation (SSST-3)*, pp.69–77, Boulder, Colorado, June 5, 2009.
104. Kei Hashimoto, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "A Bayesian approach to HMM-based speech synthesis," *2009 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2009)*, pp.4029–4032, Taipei, Taiwan, April 19–24, 2009.
105. Yi-Jian Wu, Keiichi Tokuda "Minimum generation error training by using original spectrum as reference for log spectral distortion measure," *2009 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2009)*, pp.4013–4016, Taipei, Taiwan, April 19–24, 2009.
106. Kaori Yutani, Yosuke Uto, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Voice conversion based on simultaneous modeling of spectrum and F0," *2009 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2009)*, pp.3897–3900, Taipei, Taiwan, April 19–24, 2009. Stereo-based stochastic noise compensation based on trajectory GMMs
107. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Stereo-based stochastic noise compensation based on trajectory GMMs," *2009 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2009)*, pp.4577–4580, Taipei, Taiwan, April 19–24, 2009.
108. Lu Heng, Wu Yi-Jian, Tokuda Keiichi, Dai Li-Rong, Wang Ren-Hua, "FULL covariance state duration modeling for HMM-based speech synthesis," *2009 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2009)*, pp.4033–4036, Taipei, Taiwan, April 19–24, 2009.
109. Keiichiro Oura, Yoshihiko Nankaku, Tomoki Toda, Keiichi Tokuda, Rannierry Maia, Shinsuke Sakai, Satoshi Nakamura, "Simultaneous Acoustic, Prosodic, and Phrasing Model Training for TTS Conversion Systems," *International Symposium on Chinese Spoken Language Processing (ISCSLP2008)*, SPE1.1, pp.1–4, Kunming, China, December 16–19, 2008 (Best Student Paper Award).

110. Yi-Jian Wu, Simon King and Keiichi Tokuda, "Cross-Lingual Speaker Adaptation for HMM-based Speech Synthesis," International Symposium on Chinese Spoken Language Processing (ISCSLP2008), SPE1.1, pp.9–12, Kunming, China, December 16–19, 2008.
111. Zhi-Peng Yu, Yi-Jian Wu, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Analysis of stream-dependent tying structure for HMM-based speech synthesis," International Conference on Signal Processing (ICSP'08), pp.655–658, Beijing, China, October 26–29, 2008.
112. Junichi Yamagishi, Heiga Zen, Yi-Jian Wu, Tomoki Toda, Keiichi Tokuda, "HTS-2008: Yet another evaluation of speaker adaptive HMM-based speech synthesis system," Blizzard Challenge Workshop 2008, Brisbane, Australia, September 2008 (web proceedings).
113. Ranniere Maia, Jinfu Ni, Shinsuke Sakai, Tomoki Toda, Keiichi Tokuda, Tohru Shimizu, Satoshi Nakamura, "The NICT/ATR speech synthesis system for the Blizzard Challenge 2008", Blizzard Challenge Workshop 2008, Brisbane, Australia, September 2008 (web proceedings).
114. Yoshitaka Yoshimi, Ryota Kakitsuba, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Probabilistic Answer Selection Based on Conditional Random Fields for Spoken Dialog System," Interspeech 2008, pp.215–218, Brisbane, Australia, September 22–26, 2008.
115. Yi-Jian Wu, Keiichi Tokuda, "Minimum Generation Error Training with Direct Log Spectral Distortion on LSPs for HMM-Based Speech Synthesis," Interspeech 2008, pp.577–580, Brisbane, Australia, September 22–26, 2008.
116. Sayaka Shiota, Kei Hashimoto, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Acoustic Modeling Based on Model Structure Annealing for Speech Recognition," Interspeech 2008, pp.932–935, Brisbane, Australia, September 22–26, 2008.
117. Kei Hashimoto, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Bayesian Context Clustering Using Cross Valid Prior Distribution for HMM-Based Speech Recognition," Interspeech 2008, pp.936–939, Brisbane, Australia, September 22–26, 2008.
118. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Probabilistic Feature Mapping Based on Trajectory HMMs," Interspeech 2008, pp.1068–1071, Brisbane, Australia, September 22–26, 2008.
119. Kaori Yutani, Yosuke Uto, Yoshihiko Nankaku, Tomoki Toda, Keiichi Tokuda, "Simultaneous Conversion of Duration and Spectrum Based on Statistical Models Including Time-Sequence Matching," Interspeech 2008, pp.1072–1075, Brisbane, Australia, September 22–26, 2008.
120. Tatsuya Ito, Kei Hashimoto, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Speaker Recognition Based on Variational Bayesian Method," Interspeech 2008, pp.1417–1420, Brisbane, Australia, September 22–26, 2008.
121. Simon King, Keiichi Tokuda, Heiga Zen, Junichi Yamagishi, "Unsupervised Adaptation for HMM-Based Speech Synthesis," Interspeech 2008, pp.1869–1872, Brisbane, Australia, September 22–26, 2008.

122. Yoshihiko Nankaku, Kazuhiro Nakamura, Heiga Zen, Keiichi Tokuda, Paper ID: 3947 Title: "Acoustic modeling with contextual additive structure for HMM-based speech recognition," 2008 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2008), Las Vegas, Nevada, U.S.A., March 30–April 4, 2008.
123. Junichi Yamagishi, Takashi Nose, Heiga Zen, Tomoki Toda, Keiichi Tokuda, "Performance evaluation of the speaker-independent HMM-based speech synthesis system HTS-2007 for the Blizzard Challenge 2007," 2008 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2008), Las Vegas, Nevada, U.S.A., March 30–April 4, 2008.
124. Tomoki Toda, Keiichi Tokuda, "Statistical approach to vocal tract transfer function estimation based on factor analyzed trajectory HMM," 2008 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2008), Las Vegas, Nevada, U.S.A., March 30–April 4, 2008 (accepted).
125. Ranniery Maia, Tomoki Toda, Keiichi Tokuda, Shinsuke Sakai, Satoshi Nakamura, "ON the state definition for a trainable excitation model in HMM-based speech synthesis," 2008 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2008), Las Vegas, Nevada, U.S.A., March 30–April 4, 2008 (accepted).
126. Yi-Jian Wu, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Minimum generation error criterion considering global/local variance for HMM-based speech synthesis," 2008 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2008), Las Vegas, Nevada, U.S.A., March 30–April 4, 2008 (accepted).
127. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "Model-Space MLLR for Trajectory HMMs," Interspeech 2007 - EUROSPEECH, pp.2065–2068, Antwerp, Belgium, August 27–31, 2007.
128. Ranniery Maia, Tomoki Toda, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "A trainable excitation model for HMM-based speech synthesis," Interspeech 2007 - EUROSPEECH, pp.1909–1912, Antwerp, Belgium, August 27–31, 2007.
129. Junichi Yamagishi, Heiga Zen, Tomoki Toda, Keiichi Tokuda, "Speaker-independent HMM-based speech synthesis system - HTS-2007 system for the Blizzard Challenge 2007," Proc. of Blizzard Challenge 2007, Bonn, Germany, August 25, 2007 (CD-ROM Proceedings, <http://festvox.org/blizzard/bc2007/index.html>).
130. Jinfu Ni, Toshio Hirai, Hisashi Kawai, Tomoki Toda, Keiichi Tokuda, Minoru Tsuzaki, Shinsuke Sakai, Ranniery Maia, Satoshi Nakamura, "ATRECSS – ATR English speech corpus for speech synthesis," Proc. of Blizzard Challenge 2007, Bonn, Germany, August 25, 2007 (CD-ROM Proceedings, <http://festvox.org/blizzard/bc2007/index.html>).
131. Yoshihiko Nankaku, Kenichi Nakamura, and Keiichi Tokuda, "Spectral conversion based on statistical models including time-frequency matching," Proc. of 6th ISCA Speech Synthesis Workshop, Bonn, Germany, August 22–24, 2007 (CD-ROM proceedings).
132. Shinsuke Sakai, Jinfu Ni, Ranniery Maia, Keiichi Tokuda, Minoru Tsuzaki, Tomoki Toda, Hisashi Kawai, and Satoshi Nakamura, "Communicative Speech Synthesis with XIMERA: a first step," Proc. of 6th ISCA Speech Synthesis Workshop, Bonn, Germany, August 22–24, 2007 (CD-ROM proceedings).

133. Ranniery Maia, Tomoki Toda, Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, "An excitation model for HMM-based speech synthesis based on residual modeling," Proc. of 6th ISCA Speech Synthesis Workshop, Bonn, Germany, August 22–24, 2007 (CD-ROM proceedings).
134. Junichi Yamagishi, Takao Kobayashi, Steve Renals, Simon King, Heiga Zen, Tomoki Toda, Keiichi Tokuda, "Improved average-voice-based speech synthesis using gender-mixed modeling and a parameter generation algorithm considering GV," Proc. of 6th ISCA Speech Synthesis Workshop, Bonn, Germany, August 22–24, 2007 (CD-ROM proceedings).
135. Heiga Zen, Takashi Nose, Junichi Yamagishi, Shinji Sako, Takashi Masuko, Alan W. Black, Keiichi Tokuda, "The HMM-based speech synthesis system (HTS) version 2.0," Proc. of 6th ISCA Speech Synthesis Workshop, Bonn, Germany, August 22–24, 2007 (CD-ROM proceedings).
136. Yoshihiko Nankaku, and Keiichi Tokuda, "Face recognition using hidden Markov eigen-face models," 2007 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2007), vol.2, pp.II-469–II-472, Hawaii, USA, 2007.
137. Alan W. Black, Heiga Zen, Keiichi Tokuda, "Statistical parametric speech synthesis," 2007 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2007), vol.4, pp.IV-1229–IV-1232, Hawaii, USA, 2007.
138. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, Tadashi Kitamura, "Speaker adaptation of trajectory HMMs using feature-space MLLR," Interspeech 2006 - ICSLP, pp.1141–1144, Pittsburgh, PA, Sept. 17–21, 2006.
139. Keiichi Saino, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "HMM-based singing voice synthesis system," Interspeech 2006 - ICSLP, pp.2274–2277, Pittsburgh, PA, Sept. 17–21, 2006.
140. Yosuke Uto, Yoshihiko Nankaku, Tomoki Toda, Akinobu Lee, Keiichi Tokuda, "Voice conversion based on mixtures of factor analyzers," Interspeech 2006 - ICSLP, pp.2278–2281, Pittsburgh, PA, Sept. 17–21, 2006.
141. Tomohiro Hakamata, Akinobu Lee, Yoshihiko Nankaku, Keiichi Tokuda, "Reducing computation on Parallel decoding using frame-wise confidence scores," Interspeech 2006 - ICSLP, pp.1638–1641, Pittsburgh, PA, Sept. 17–21, 2006.
142. Heiga Zen, Tomoki Toda, Keiichi Tokuda, "The Nitech-NAIST HMM-based speech synthesis system for the Blizzard Challenge 2006," Blizzard Challenge 2006 Workshop, Pittsburgh, PA, 2006 (<http://festvox.org/blizzard/blizzard2006.html>).
143. Tomoki Toda, Hisashi Kawai, Toshio Hirai, Jinfu Ni, Nobuyuki Nishizawa, Junichi Yamagishi, Minoru Tsuzaki, Keiichi Tokuda, Satoshi Nakamura, "Developing a Test Bed of English Text-to-Speech System XIMERA for the Blizzard Challenge 2006," Blizzard Challenge 2006 Workshop, Pittsburgh, PA, 2006 (<http://festvox.org/blizzard/blizzard2006.html>).

144. Heiga Zen, Yoshihiko Nankaku, Keiichi Tokuda, Tadashi Kitamura, "Estimating trajectory HMM parameters using Monte Carlo EM with Gibbs sampler," 2006 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2006), vol.1, pp.I-1173–I-1176, Toulouse, France, May 14-19, 2006.
145. Keiichiro Oura, Heiga Zen, Yoshihiko Nankaku, Akinobu Lee, Keiichi Tokuda, "Hidden semi-Markov model based speech recognition system using weighted finite-state transducer," 2006 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2006), vol.1, pp.I-33–I-36, Toulouse, France, May 14-19, 2006.
146. Kenichi Nakamura, Tomoki Toda, Yoshihiko Nankaku, Keiichi Tokuda, "On the use of phonetic information for mapping from articulatory movements to vocal tract spectrum," 2006 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2006), vol.1, pp.I-93–I-96, Toulouse, France, May 14-19, 2006.
147. Daisuke Kurata, Yoshihiko Nankaku, Keiichi Tokuda, Tadashi Kitamura, Zoubin Ghahramani, "Face recognition based on separable lattice HMMs," 2006 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2006), vol.5, pp.V-737–V-740, Toulouse, France, May 14-19, 2006 (Student Paper Award).
148. Alan W. Black, Keiichi Tokuda, "The Blizzard Challenge – 2005: Evaluating corpus-based speech synthesis on common datasets," INTERSPEECH 2005 - EUROSPEECH, pp.77–80, Lisbon, Portugal, September 4–8, 2005.
149. Maria João Barros, Ranniery Maia, Keiichi Tokuda, Fernando Gil Resende, Diamantino Freitas, "HMM-based European Portuguese TTS system," INTERSPEECH 2005 - EUROSPEECH, pp.2581–2584, Lisbon, Portugal, September 4–8, 2005.
150. Tomoki Toda, Keiichi Tokuda, "Speech parameter generation algorithm considering global variance for HMM-based speech synthesis," INTERSPEECH 2005 - EUROSPEECH, pp.2801–2804, Lisbon, Portugal, September 4–8, 2005.
151. Yusuke Kida, Hiroyoshi Yamamoto, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Minimum classification error interactive training for speaker identification," 2005 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2005), pp.I-641–I-644, Philadelphia, USA, March 18-23, 2005.
152. Amaro Azevedo de Lima, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Sparse KPCA for feature extraction in speech recognition," 2005 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2005), pp.I-353–I-356, Philadelphia, USA, March 18-23, 2005.
153. Tomoki Toda, Alan W. Black, Keiichi Tokuda, "Spectral conversion based on maximum likelihood estimation considering global variance of converted parameter," 2005 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2005), pp.I-9–I-12, Philadelphia, USA, March 18-23, 2005.
154. Keiichi Tokuda, Heiga Zen, and Tadashi Kitamura, "Reformulating the HMM as a Trajectory Model," Workshop on Statistical modeling Approach for Speech Recognition (Beyond HMM), Kyoto, Japan, Dec. 20, 2004.

155. Shunsuke Kataoka, Nobuaki Mizutani, Keiichi Tokuda, Tadashi Kitamura, "Decision-tree backing-off in HMM-based speech synthesis," International Conference on Spoken Language Processing (INTERSPEECH2000-ICSLP2000), vol.2, pp.II-1205–II-1208, Jeju, Korea, Oct. 4–8, 2004.
156. Ryosuke Tsuzuki, Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, Murtaza Bulut, Shrikanth S. Narayanan, "Constructing emotional speech synthesizers with limited speech database," International Conference on Spoken Language Processing (INTERSPEECH2004-ICSLP2004), vol.2, pp.II-1185–II-1188, Jeju, Korea, Oct. 4–8, 2004.
157. Tomoki Toda, Alan W. Black, Keiichi Tokud "Acoustic-to-articulatory inversion mapping with Gaussian mixture model," International Conference on Spoken Language Processing (INTERSPEECH2004-ICSLP2004), vol.2, pp.II-1129–II-1132, Jeju, Korea, Oct. 4–8, 2004.
158. Heiga Zen, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "Hidden semi-Markov model based speech synthesis," International Conference on Spoken Language Processing (INTERSPEECH2004-ICSLP2004), vol.2 pp.II-1393–III-386, Jeju, Korea, Oct. 4–8, 2004.
159. Yohei Itaya, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Deterministic annealing EM algorithm in parameter estimation for acoustic model," International Conference on Spoken Language Processing (INTERSPEECH2004-ICSLP2004), vol.1, pp.I-437—I-440, Jeju, Korea, Oct. 4–8, 2004.
160. T. Nitta, S. Sagayama, Y. Yamashita, T. Kawahara, S. Morishima, S. Nakamura, A. Yamada, K. Ito, M. Kai, A. Li, M. Mimura, K. Hirose, T. Kobayashi, K. Tokuda, N. Minematsu, Y. Den, T. Utsuro, T. Yotsukura, H. Shimodaira, M. Araki, T. Nishimoto, N. Kawaguchi, H. Banno, K. Katsurada, "Activities of Interactive Speech Technology Consortium (ISTC) Targeting Open Software Development for MMI Systems," 13th IEEE International Workshop on Robot and Human Interactive Communication (RO-MAN 2004), Kurashiki, Okayama, Japan, September 20-22, 2004 (CD-ROM proceedings).
161. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "An introduction of trajectory model into HMM-based speech synthesis," Proc. of 5th ISCA Speech Synthesis Workshop, Pittsburgh, U.S.A., June 2004 (CD-ROM proceedings).
162. Tomoki Toda, Alan W. Black, and Keiichi Tokuda, "Mapping from articulatory movements to vocal tract spectrum with gaussian mixture model for articulatory speech synthesis," Proc. of 5th ISCA Speech Synthesis Workshop, Pittsburgh, U.S.A., June 2004 (CD-ROM proceedings).
163. Hisashi Kawai, Tomoki Toda, Jinfu Ni, Minoru Tsuzaki, and Keiichi Tokuda, "Ximera: A New TTS from ATR Based on Corpus-Based Technologies," ISCA 5th Speech Synthesis Workshop, pp.179-184, Pittsburgh, U.S.A., June 2004.
164. Hiroyoshi Yamamoto, Yoshihoko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Parameter sharing and minimum classification error training of mixtures of factor analyzers for speaker identification," 2004 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2004), vol.1, pp.29–32, Montreal, Canada, May 17-21, 2004.

165. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "A Viterbi algorithm for a trajectory model derived from HMM with explicit relationship between static and dynamic features," Proceedings of 2004 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP 2004), vol.1, pp.837–840, Montreal, Canada, May 17-21, 2004.
166. Takahiro Hoshiya, Heiga Zen, Shinji Sako, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "An HMM-based approach to speaker-dependent 100 bit/s speech coding," Special Workshop in MAUI (SWIM), Lectures by Masters in Speech Processing, Maui, Hawaii, Jan. 12-14, 2004 (invited session).
167. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "Decision Tree-based Simultaneous Clustering of Phonetic Contexts, Dimensions, and State Positions for Acoustic Modeling," Proceedings of European Conference on Speech Communication and Technology, pp.3189–3192, Geneva, Switzerland, Sep. 1–4, 2003.
168. Amaro Azevedo de Lima, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "On the Use of Kernel PCA for Feature Extraction in Speech Recognition," Proceedings of European Conference on Speech Communication and Technology, pp.2625–2628, Geneva, Switzerland, Sep. 1–4, 2003.
169. Ranniery da Silva Maia, Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, Fernando Gil Vianna Resende Junior, "Towards the development of a Brazilian Portuguese text-to-speech system based on HMM," Proceedings of European Conference on Speech Communication and Technology, pp.2465–2468, Geneva, Switzerland, Sep. 1–4, 2003.
170. Keiichi Tokuda, Heiga Zen, Tadashi Kitamura, "Trajectory modeling based on HMMs with the explicit relationship between static and dynamic features," Proceedings of European Conference on Speech Communication and Technology, pp.865–868, Geneva, Switzerland, Sep. 1–4, 2003.
171. Junichi Yamagishi, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "A training method for average voice model based on shared decision tree context clustering and speaker adaptive training," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), vol.1, pp.716–719, Hong Kong, April 6–10, 2003.
172. Hiroyuki Suzuki, Heiga Zen, Yoshihiko Nankaku, Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Speech recognition using voice-characteristic-dependent acoustic models," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), vol.1, pp.740–743, Hong Kong, April 6–10, 2003.
173. Takahiro Hoshiya, Shinji Sako, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, Heiga Zen, "Improving the performance of HMM-based very low bitrate speech coding," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), vol.1, pp.800–803, Hong Kong, April 6–10, 2003.
174. Keiichi Tokuda, Zen Heiga, Alan W. Black, "An HMM-based speech synthesis system applied to English," 2002 IEEE Speech Synthesis Workshop, Santa Monica, California, Sep. 11–13, 2002 (CD-ROM proceedings).

175. Shin-ichi Kawamoto, Hiroshi Shimodaira, Tsuneo Nitta, Takuya Nishimoto, Satoshi Nakamura, Katsunobu Itou, Shigeo Morishima, Tatsuo Yotsukura, Atsuhiko Kai, Akinobu Lee, Yoichi Yamashita, Takao Kobayashi, Keiichi Tokuda, Keikichi Hirose, Nobuaki Minematsu, Atsushi Yamada, Yasuharu Den, Takehito Utsuro, Shigeki Sagayama, "Open-source Software for Developing Anthropomorphic Spoken Dialog Agents," International Workshop on LIFELIKE ANIMATED AGENTS Tools, Affective Functions, and Applications, pp.64–69, Tokyo, Japan, August 19, 2002.
176. Junichi Yamagishi, Masatsune Tamura, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "A context clustering technique for average voice model in HMM-based speech synthesis," 8th International Conference on Spoken Language Processing, pp.133–136, Sep. 16–20, Denver, Colorado, 2002.
177. Kengo Shichiri, Atsushi Sawabe, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "Eigenvoices for HMM-based speech synthesis," 8th International Conference on Spoken Language Processing, pp.1269–1272, Sep. 16–20, Denver, Colorado, 2002.
178. Heiga Zen, Keiichi Tokuda, Tadashi Kitamura, "Decision tree distribution tying based on a dimensional split technique," 8th International Conference on Spoken Language Processing, pp.1257–1260, Sep. 16–20, Denver, Colorado, 2002.
179. Shin-ichi Kawamoto, Hiroshi Shimodaira, Shigeki Sagayama, Tsuneo Nitta, Takuya Nishimoto, Satoshi Nakamura, Katsunobu Itou, Shigeo Morishima, Tatsuo Yotsukura, Akinobu Kai, Akinobu Lee, Yoichi Yamashita, Takao Kobayashi, Keiichi Tokuda, Keikichi Hirose, Nobuaki Minematsu, Atsushi Yamada, Yasuharu Den, Takehito Utsuro, "Developments of anthropomorphic dialog agent: A plan and development and its significance," International Workshop on Information Presentation and Natural Multimodal Dialogue, Verona, Italy, 14–15 Dec. 2001.
180. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "Mixed excitation for HMM-based speech Synthesis," Proceedings of European Conference on Speech Communication and Technology, vol.3, pp.2263-2266, Aalborg, Denmark, Sep. 3–7, 2001.
181. Takayuki Satoh, Takashi Masuko, Takao Kobayashi, Keiichi Tokuda, "A robust speaker verification system against imposture using an HMM-based speech synthesis system," Proceedings of European Conference on Speech Communication and Technology, vol.2, pp.759–762, Aalborg, Denmark, Sep. 3–7, 2001.
182. Masatsune Tamura, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi, "Text-to-speech synthesis with arbitrary speaker's voice from average voice," Proceedings of European Conference on Speech Communication and Technology, vol.1, pp.345-348, Aalborg, Denmark, vol.1 pp.345–348, Sep. 3–7, 2001.
183. Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Minimum classification error training for speaker identification using Gaussian mixture models based on multi-space probability distribution," Proceedings of European Conference on Speech Communication and Technology, Aalborg, Denmark, vol.4, pp.2347–2350, Sep. 3–7, 2001.

184. Masatsune Tamura, Takashi Masuko, Keiichi Tokuda and Takao Kobayashi, "Adaptation of pitch and spectrum for HMM-based speech synthesis using MLLR," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), vol.2, pp.805–808, Salt Lake City, Utah, USA, May 7–11, 2001.
185. Chiyomi Miyajima, Yosuke Hattori, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, and Tadashi Kitamura, "Speaker identification using Gaussian mixture models based on multi-space probability distribution," Proceedings of International Conference on Acoustics, Speech, and Signal Processing (ICASSP), Salt Lake City, Utah, USA, pp.433–436, May 7–11, 2001.
186. Junibakti Sanubari, Keiichi Tokuda, "Fast convergence transversal adaptive filtering algorithm for impulsive environment based on t-distribution," 2001 IEEE International Symposium on Circuits and Systems, Sydney, Australia, May 6 - 9, 2001.
187. Takahiro Nakanishi, Keiichi Tokuda, Tadashi Kitamura, "Simultaneous determination of model-order and frame partitioning for time series analysis based on MDL criterion," International Symposium on Information Theory and Its Applications (ISITA-2000), Honolulu, Hawaii, November 5-8, 2000 (accepted).
188. Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Audio-visual speech recognition based on minimum classification error discriminative training," 2000 IEEE International Workshop on Neural Networks for Signal Processing, Sydney, Australia, pp.3–12, 11–13 Dec. 2000 (invited session).
189. Shinji Sako, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "HMM-based text-to-audio-visual speech synthesis," International Conference on Spoken Language Processing (ICSLP2000/INTERSPEECH2000), vol.III, pp.25–28, Beijing, China, Oct. 16–20, 2000 (invited session).
190. Chiyomi Miyajima, Keiichi Tokuda, Tadashi Kitamura, "Audio-visual speech recognition using MCE-based HMMs and model-dependent stream weights," International Conference on Spoken Language Processing (ICSLP2000/INTERSPEECH2000), vol.II, pp.1023–1026, Beijing, China, Oct. 16–20, 2000.
191. Takashi Masuko, Keiichi Tokuda and Takao Kobayashi, "Imposture using synthetic speech against speaker verification based on spectrum and pitch," International Conference on Spoken Language Processing (ICSLP2000/INTERSPEECH2000), vol.III, pp.302–305, Beijing, China, Oct. 16–20, 2000.
192. Toru Takahashi, Keiichi Tokuda, Takao Kobayashi, Tadashi Kitamura, "Vector quantization of mel-cepstral coefficients based on a statistical measure," IEEE International Symposium on Intelligent Signal Processing and Communication Systems, Honolulu, Hawaii, pp.692–695, Nov. 5-8, 2000.
193. Junibakti Sanubari, Keiichi Tokuda, "Robust estimation of an AR multi-channel model by using t-distribution assumption," Proc. The 10th European Signal Processing Conference (EUSIPCO-2000), vol.3, pp.1783–1786, Tampere, Finland, 4–8, Sep. 2000.
194. Yoshihiko Nankaku, Keiichi Tokuda, Takao Kobayashi, Tadashi Kitamura, "Normalized training for HMM-based visual speech recognition," Proc. of IEEE International Conference on Image Processing, Vancouver, Canada, vol.3, pp.234–237, Sep. 2000.

195. Keiichi Tokuda, Takayoshi Yoshimura, Takashi Masuko, Takao Kobayashi, Tadashi Kitamura, "Speech parameter generation algorithms for HMM-based speech synthesis," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, Istanbul, Turkey, vol.3, pp.1315–1318, June 2000.
196. Junibakti Sanubari, Keiichi Tokuda, "Image modeling using two dimensional exponential systems," IEEE International Conference on Image Processing, Kobe, Japan, pp.28AP4.8.1–28AP4.8.4, Oct. 1999.
197. Junibakti Sanubari, Keiichi Tokuda, "Two dimensional adaptive filter based on t-distribution assumption and full-plane support," Proceedings of Thirty-Third Asilomar Conference on Signal, Systems, and Computers, Pacific Grove, California, Oct. 24–27, pp.815–819, 1999.
198. Oscar Vanegas, Keiichi Tokuda, Tadashi Kitamura, "Location normalization of HMM-based lip reading: Experiments for the M2VTS Database," IEEE International Conference on Image Processing, Kobe, Japan, Oct. 1999 (accepted).
199. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Simultaneous modeling of spectrum, pitch and duration in HMM-based speech synthesis," Proceedings of European Conference on Speech Communication and Technology, Budapest, Hungary, vol.5, pp.2347–2350, Sep. 1999.
200. Yoshihiko Nankaku, Keiichi Tokuda and Tadashi Kitamura, "Intensity- and location-normalized training for HMM-based visual speech recognition," Proceedings of European Conference on Speech Communication and Technology, Budapest, Hungary, vol.3, pp.1287–1290, Sep. 1999.
201. Takashi Masuko, Takafumi Hitotsumatsu, Keiichi Tokuda and Takao Kobayashi, "On the security of HMM-based speaker verification systems against imposture using synthetic speech," Proceedings of European Conference on Speech Communication and Technology, Budapest, Hungary, vol.3, pp.1223–1226, Sep. 1999.
202. Junibakti Sanubari, Keiichi Tokuda, "LMS-like two dimensional adaptive filter with t-distribution assumption and non-symmetrical half plane support," Proceedings of IEEE-EURASIP Workshop on Nonlinear and Image Processing, Antalya, Turkey, pp.419–423, June 1999.
203. Keiichi Tokuda, Takashi Masuko, Noboru Miyazaki and Takao Kobayashi, "Hidden Markov models based on multi-space probability distribution for pitch pattern modeling," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.229–232, Phoenix, USA, Mar. 1999.
204. Masatsune Tamura, Takashi Masuko, Takao Kobayashi and Keiichi Tokuda, "Visual speech synthesis based on parameter generation from HMM: speech-driven and text-and-speech-driven approach," Proc. International Conference of Auditory-Visual Speech Processing, pp.219–224, Terrigal, Australia, Dec. 1998.
205. Kazuhito Koishida, Goh Hirabayashi, Keiichi Tokuda, and Takao Kobayashi, "A 16kbit/s wideband CELP coder using mel-generalized cepstral analysis and its subjective evaluation," Proc. of International Conference on Spoken Language Processing (ICLSP-98), vol.6., pp.2583–2586, Sydney, Australia, Nov.–Dec., 1998.

206. Oscar Vanegas, Akiji Tanaka, Keiichi Tokuda, Tadashi Kitamura, "HMM-based visual speech recognition using intensity and location normalization," Proc. of International Conference on Spoken Language Processing (ICLSP-98), vol.2, pp.789–792, Sydney, Australia, Nov.–Dec., 1998.
207. Takayoshi Yoshimura, Keiichi Tokuda, Takashi Masuko, Takao Kobayashi and Tadashi Kitamura, "Duration modeling for HMM-based speech synthesis," Proc. of International Conference on Spoken Language Processing (ICLSP-98), vol.2, Sydney, Australia, pp.29–32, Nov.–Dec., 1998.
208. Takashi Masuko, Takao Kobayashi and Keiichi Tokuda, "A very low bit rate speech coder using HMM with speaker adaptation," Proc. of International Conference on Spoken Language Processing (ICLSP-98), vol.2, pp.507–510, Sydney, Australia, Nov.–Dec., 1998.
209. Masatsune Tamura, Takashi Masuko, Keiichi Tokuda and Takao Kobayashi, "Speaker adaptation for HMM-based speech synthesis system using MLLR," Proc. ESCA/COCOSDA Workshop on Speech Synthesis, pp.273–276, Blue Mountains, Australia, Nov. 1998.
210. Oscar Vanegas, Akiji Tanaka, Keiichi Tokuda, Tadashi Kitamura, "Intensity/location normalization for automatic lipreading," Proc. International Conference on Signal Processing, vol.2, pp.920–923, Oct. 1998.
211. Junibakti Sanubari, Keiichi Tokuda, "Recursive two dimensional spectral estimation based on an AR model excited by a t-distribution process using AR decomposition approach," Proceedings of IEEE Asia Pacific Conference on Circuits and Systems, Chiang Mai, Thailand, pp.447–450, Nov. 1998.
212. Junibakti Sanubari, Keiichi Tokuda, "Adaptive two dimensional filter based on an AR model excited by a t-distribution process," Proceedings of IESTED International Conference on Signal and Image Processing, Las Vegas, Nevada-USA, pp.679–683, Oct. 1998.
213. Junibakti Sanubari, Keiichi Tokuda, "Adaptive spectral estimation based on an exponential model," Proceedings of IEEE International Conference on Circuits and Systems, Monterey, California, pp.TAA13-11.1–TAA13-11.4, May 1998.
214. Keiichi Tokuda, Takashi Masuko, Jun Hiroi, Takao Kobayashi, and Tadashi Kitamura, "A very low bit rate speech coder using HMM-based speech recognition/synthesis techniques," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.2, pp.609–612, May 1998.
215. Kazuhito Koishida, Goh Hirabayashi, Keiichi Tokuda, and Takao Kobayashi, "A wide-band CELP speech coder at 16 kbit/s based on mel-generalized cepstral analysis," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.161–164, Seattle, USA, May 1998.
216. Takashi Masuko, Takao Kobayashi, Masatsune Tamura, Jun Masubuchi, and Keiichi Tokuda, "Text-to-visual speech synthesis based on parameter generation from HMM," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.6, pp.3745–3748, Seattle, USA, May 1998.

217. Junibakti Sanubari and Keiichi Tokuda, "Non stationary spectral estimation based on robust time varying AR model excited by a t-distribution process," IEEE Region Ten Annual Conference on Speech and Image Technologies for Computing and Telecommunications, Brisbane, Australia, pp.51–54, Dec. 1997.
218. Takao Kobayashi, Takashi Masuko and Keiichi Tokuda, "HMM compensation for noisy speech recognition based on cepstral parameter generation," Proceedings of European Conference on Speech Communication and Technology, vol.3, pp.1583–1586, Rhodes, Greece, Sep. 1997.
219. Takayoshi Yoshimura, Takashi Masuko, Keiichi Tokuda, Takao Kobayashi and Tadashi Kitamura, "Speaker interpolation in HMM-based speech synthesis system," Proceedings of European Conference on Speech Communication and Technology, vol.5, pp.2523–2526, Rhodes, Greece, Sep. 1997.
220. Kazuhito Koishida, Keiichi Tokuda, Takashi Masuko and Takao Kobayashi, "Spectral quantization using statistics of static and dynamic features," Proc. of 1997 IEEE Workshop on Speech Coding for Telecommunications, pp.19–20, Pocono Manor, USA, Sep. 1997.
221. Chiyomi Miyajima, Yoshitsuna Sugiura, Keiichi Tokuda and Tadashi Kitamura, "Discrete or tied-mixture HMM based of self-organizing feature map for robust probability estimation," Proceedings of International Conference on Speech Processing, vol.2, pp.529–532, Seoul, Korea, Aug. 1997.
222. Kazuhito Koishida, Keiichi Tokuda, Takashi Masuko and Takao Kobayashi, "Vector quantization of speech spectral parameters using statistics of dynamic features," Proceedings of International Conference on Speech Processing, vol.1, pp.247–252, Seoul, Korea, Aug. 1997.
223. Junibakti Sanubari, Keiichi Tokuda, "Robust spectral estimation based on an AR model excited by a t-distribution process by using QR decomposition algorithm," Proceedings of IEEE International Conference on Circuits and Systems, pp.2497–2500, June 1997.
224. Fernando G. Resende, Keiichi Tokuda and Mineo Kaneko, "Multi-band decomposition of the linear prediction error applied to the least-mean squares method with fixed and variable step-sizes," Proceedings of IEEE International Conference on Circuits and Systems, pp.2176–2179, June 1997.
225. Takashi Masuko, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Voice characteristics conversion for HMM-based speech synthesis system," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.3, pp.1611–1614, Munich, Germany, Apr. 1997.
226. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Efficient encoding of mel-generalized cepstrum for CELP coders," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.2, pp.1355–1358, Munich, Germany, Apr. 1997.
227. Yasuichi Hamano, Keiichi Tokuda and Mineo Kaneko, "Image restoration based on estimation of fractal structure," IEEE Region Ten Conference (Digital Signal Processing Applications), pp.311–316, Nov. 1996.

228. Fernando G. Resende, Keiichi Tokuda and Mineo Kaneko, "RLS algorithms for adaptive AR spectrum analysis based on multi-band decomposition of the linear prediction error," IEEE Region Ten Conference (Digital Signal Processing Applications), pp.541-546, Nov. 1996.
229. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "CELP coding system based on mel-generalized cepstral analysis," Proceedings of International Conference on Spoken Language Processing, Philadelphia, USA, vol.1, pp.314-317, Oct. 1996.
230. Junibakti Sanubari, Keiichi Tokuda and Mahoki Onoda, "The application of t-distribution assumption for robust spectral estimation," Proc. of the IASTED International Conferences on Signal, Image Processing and Application, SIPA-96, Annecy, France, pp.217-220, June 1996.
231. Junibakti Sanubari, Keiichi Tokuda, Mahoki Onoda, "Robust two-dimensional spectral estimation based on an AR model excited by a t-distribution process," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.5, pp.2998-3001, Atlanta, USA, May 1996.
232. Takashi Masuko, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Speech synthesis from HMMs using dynamic features," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.389-392, May 1996.
233. Fernando G. Resende, Keiichi Tokuda and Mineo Kaneko, "A fast algorithm for adaptive AR spectral estimation based on multi-scale decomposition of linear prediction error," Proceedings of Midwest Symposium on Circuits and Systems, pp.119-122, Aug. 1995.
234. Keiichi Tokuda, Takashi Masuko, Tetsuya Yamada, Takao Kobayashi and Satoshi Imai, "An algorithm for speech parameter generation from continuous mixture HMMs with dynamic features," Proceedings of European Conference on Speech Communication and Technology, vol.1, pp.757-760, Sep. 1995.
235. Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Speech parameter generation from HMM using dynamic features," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.660-663, May 1995.
236. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "CELP coding based on mel-cepstral analysis," Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.33-36, May 1995.
237. Kazuhito Koishida, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, "Speech coding based on adaptive mel-cepstral analysis for noisy channels," Proceedings of International Conference on Spoken Language Processing, vol.4, pp.2087-2090, Sep. 1994.
238. Keiichi Tokuda, Takao Kobayashi, Takashi Masuko and Satoshi Imai, "Mel-generalized cepstral analysis —a unified approach to speech spectral estimation," Proceedings of International Conference on Spoken Language Processing, vol.3, pp.1043-1046, Sep. 1994.
239. Fernando G. Resende, Keiichi Tokuda and Mineo Kaneko, "AR spectrum estimation based on wavelet representation," Proceedings of IEEE International Conference on Circuits and Systems, vol.2, pp.625-628, June 1994.

240. Junibakti Sanubari, Keiichi Tokuda, Mahoki Onoda, “Robust recursive spectral estimation based on AR model excited by a t-distribution process,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.3, pp.497–500, Apr. 1994.
241. Keiichi Tokuda, Hidetoshi Matsumura, Takao Kobayashi and Satoshi Imai, “Speech coding based on adaptive mel-cepstral analysis,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.197–200, Apr. 1994.
242. Junibakti Sanubari, Keiichi Tokuda and Mahoki Onoda, “Spectral estimation based on AR model excited by t-distribution process,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.5, pp.521–524, Mar. 1992.
243. Takao Kobayashi, Kazuyoshi Fukushi, Keiichi Tokuda and Satoshi Imai, “Design of stable two-dimensional IIR digital filters with arbitrary magnitude function,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.5, pp.93–96, Mar. 1992.
244. Toshiaki Fukada, Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, “An adaptive algorithm for mel-cepstral analysis of speech,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, vol.1, pp.137–140, Mar. 1992.
245. Keiichi Tokuda, Takao Kobayashi and Satoshi Imai, “Generalized cepstral analysis of speech —unified approach to LPC and cepstral method,” Proceedings of International Conference on Spoken Language Processing, pp.37–40, Nov. 1990.
246. Keiichi Tokuda, Takao Kobayashi, Shoji Shiimoto and Satoshi Imai, “Adaptive filtering based on cepstral representation —adaptive cepstral analysis of speech,” Proceedings of IEEE International Conference on Acoustics, Speech, and Signal Processing, pp.377–380, Apr. 1990.

5 SOFTWARE ARTIFACTS

The HMM-based Speech Synthesis System (HTS) <http://hts.sp.nitech.ac.jp/>

The Speech Signal Processing Toolkit (SPTK) <http://sp-tk.sourceforge.net/>

Multimodal Speech Database (M2TINIT) <http://m2tinit.sp.nitech.ac.jp/>

The Galatea Toolkit <http://hil.t.u-tokyo.ac.jp/~galatea/>

hts_engine API <http://hts-engine.sourceforge.net/>

HMM-based Japanese TTS system (Open JTalk) <http://open-jtalk.sourceforge.net/>

A toolkit for building voice interaction systems (MMDAgent) <http://www.mmdagent.jp/>

HMM-based singing voice synthesis system (Sinsy) <http://sinsy.sourceforge.net/>

6 EXTERNAL TEACHING EXPERIENCE

Nov. 29, 2008

Invited Lecturer, Meijo University, Japan

Aug. 8, 2003

LTI Seminar, Language Technologies Institute, Carnegie Mellon University

Mar. 12–13, 2002

International Lecture Series on Signal Processing, Tamkang University Lecture Series of EECS, Tamkang University, R.O.C.

Feb. 8, 2002

LTI Seminar, Language Technologies Institute, Carnegie Mellon University

Dec. 4, 2001

Speech and Signal Processing Seminar, Department of Electrical Engineering, University of Washington

April 2001–September 2001

Invited Lecturer, Nagoya University, Japan

November 1997–March 1998

Invited Lecturer, Tokyo Institute of Technology, Japan

7 INVITED TALKS

AT INTERNATIONAL CONFERENCES

Keiichi Tokuda, “Thirty Years of Progress in Speech Synthesis: A Personal Perspective on the Past, Present, and Future,” Symposium on Speech & Behavior Informatics, December 2025.

Keiichi Tokuda, “Statistical Approach to Speech Synthesis: Past, Present and Future,” Interspeech 2019, September, 2019.

Keiichi Tokuda, “Flexible speech synthesis based on hidden Markov models,” Asia-Pacific Signal and Information Processing Association Annual Summit and Conference 2013 (APSIPA ASC 2013), Kaohsiung, Taiwan, October 29–November 1, 2013.

Keiichi Tokuda, “Speech Synthesis as A Statistical Machine Learning Problem,” IEEE 2011 Automatic Speech Recognition and Understanding Workshop (ASRU 2011), Hawaii, USA, December 11–15, 2011.

Keiichi Tokuda, “Speech Synthesis as A Machine Learning Problem,” The International Committee for the Co-ordination and Standardization of Speech Databases and Assessment Techniques (Oriental COCOSDA 2010), Kathmandu, Nepal, Nov. 24–25, 2010.

Keiichi Tokuda, Heiga Zen, “Fundamentals and recent advances in HMM-based speech synthesis,” Interspeech 2009, Tutorial, Brighton, U.K., September 6, 2009.

Keiichi Tokuda, “An HMM-based approach to flexible speech synthesis,” The 5th International Symposium on Chinese Spoken Language Processing (ISCSLP 2006), Kent Ridge, Singapore, December 13–16, 2006 (tutorial).

Keiichi Tokuda, “Hidden Markov model-based speech synthesis as a tool for constructing communicative spoken dialog systems,” Proc. of The 4th Joint Meeting of The Acoustical Society of America and The Acoustical Society of Japan, Honolulu, Hawaii, 28 November–2 December 2006 (in J. Acoust. Soc. Am., vol.120, no.5, Part.2, p.3006, November 2006) (invited talk).

Keiichi Tokuda, “The trajectory HMM: Reformulating the HMM as a trajectory model,” Trajectory Models for Speech Processing, Supported by EPSRC (The Engineering and Physical Sciences Research Council, U.K.), Edinburgh, U.K., August 31, 2005 (keynote).

More than 16 talks were given at domestic conferences

AT COMPANIES, UNIVERSITIES, AND OTHER WORKSHOPS

The University of Cambridge, CUED Speech Group Seminars, Cambridge, U.K., “Human-like singing and talking machines: flexible speech synthesis in karaoke, anime, smart phones, video games, digital signage, TV and radio programs, etc.,” February 6, 2015.

The University of Sheffield, Speech and Hearing Seminars, U.K., “Human-like singing and talking machines,” January 13, 2015.

Massachusetts Institute of Technology, Human Language Technology Distinguished Lecture Series, MA, U.S.A., “Human-like Singing and Talking Machines: Flexible Speech Synthesis in Karaoke, Anime, Smart Phones, Video Games, Digital Signage, TV and Radio Programs,” December 4, 2014.

5th EU-Japan Symposium in ICT Research and Innovation, Brussels, Belgium, “Statistical approach to flexible speech synthesis –toward human-like talking machines,” October 16–17, 2014.

University of Science and Technology of China, Talk, “Fundamentals and recent advances in HMM-based speech synthesis,” October 30, 2008.

University of Science and Technology of China, “HMM-Based Speech Synthesis toward human-like talking machines,” March 6, 2007.

Microsoft Research Asia, Beijing, China, “An HMM-based approach to automatic voice building,” March 28, 2006.

University of Cambridge, Department of Engineering, Speech Vision and Robotics Group, Cambridge, U.K., “HMM-based speech synthesis —toward Human-like Talking Machines,” Apr. 26, 2002.

Carnegie Mellon University, PA, U.S.A., Talk at Sphinx Lunch, “HMM-based speech synthesis —toward Human-like Talking Machines,” Apr. 18, 2002.

Cepstral, PA, U.S.A., , “HMM-based speech synthesis —toward Human-like Talking Machines,” Apr. 10, 2002.

Lucent Technologies, Bell Labs, NJ, U.S.A., “HMM-based speech synthesis —toward Human-like Talking Machines,” Apr. 2, 2002.

AT&T Labs Research, “HMM-based speech synthesis —toward Human-like Talking Machines,” Mar. 26, 2002.

Nuance Communications, CA, U.S.A., “HMM-based speech synthesis —toward human-like talking machines,” Feb. 28, 2002.

Universidade Federal do Rio de Janeiro, Signal Processing Laboratory, Rio de Janeiro, Brazil, “HMM-based speech synthesis —toward human-like talking machines,” Feb. 19, 2002.

IBM Research, Watson Research Center, NY, U.S.A., “HMM-based speech synthesis —toward human-like talking machines,” Feb. 12, 2002.

Oregon Graduate Institute, Center for Spoken Language Understanding, OR, U.S.A., , “HMM-based speech synthesis —toward human-like talking machines,” Dec. 6, 2001.

Microsoft Research, WA, U.S.A., “HMM-based speech synthesis —toward Human-like Talking Machines,” Nov. 30, 2001.

More than 14 talks were given at domestic companies and universities.

8 PROFESSIONAL ACTIVITIES

CONFERENCE AND WORKSHOP COMMITTEES

The Blizzard Challenge Workshop Co-organizer, Vienna, Austria, September 23, 2019.

Interspeech 2018 satellite workshop “The Blizzard Challenge Workshop 2018,” Co-organizer, Hyderabad, India, September 8, 2018.

The Blizzard Challenge Workshop Co-organizer, Stockholm, Sweden, August 25, 2017 (satellite event of Interspeech 2017).

The Blizzard Challenge Workshop Co-organizer, California, USA, September 16, 2016 (satellite event of 9th ISCA Speech Synthesis Workshop and Interspeech 2016).

9th ISCA Speech Synthesis Workshop Program Committee Member, California, USA, September 13–15, 2016.

Interspeech 2015 satellite workshop “The Blizzard Challenge 2015” Co-organizer, Berlin, Germany, September 11, 2015.

Interspeech 2014 satellite workshop “The Blizzard Challenge 2014: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Singapore, Singapore, September 19, 2014.

The Blizzard Challenge Workshop Co-organizer, Barcelona, Spain, September 3, 2013 (satellite event of 8th ISCA Speech Synthesis Workshop).

8th ISCA Speech Synthesis Workshop Advisory Committee Member, Barcelona, Spain, August 31–September 2, 2013.

Interspeech 2012 satellite workshop “The Blizzard Challenge 2008: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Portland, Oregon, USA, September 14, 2012.

Interspeech 2011 satellite workshop “The Blizzard Challenge 2008: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Turin, Italy, September 2, 2011.

Interspeech 2010 Special Session Chair, Makuhari, Japan, September 26–30, 2010.

The Blizzard Challenge Workshop Co-organizer, Kyoto, Japan, September 25, 2010 (satellite event of 7th ISCA Speech Synthesis Workshop).

7th ISCA Speech Synthesis Workshop Co-Chair, Kyoto, Japan, September 22–24, 2010.

Interspeech 2009 satellite workshop “The Blizzard Challenge 2008: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Edinburgh, U.K., September 4, 2009.

ACL-08 Area Chair, Columbus, Ohio, USA, June 15–20, 2008.

Interspeech 2008 satellite workshop “The Blizzard Challenge 2008: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Brisbane, Australia, September 21, 2008.

IEEE ICASSP 2007 “Statistical Parametric Speech Synthesis,” Special Session Co-organizer, Hawaii, USA, April 15–20, 2007.

The Blizzard Challenge Workshop Co-organizer, Bonn, Germany, August 25, 2007 (satellite event of 6th ISCA Speech Synthesis Workshop).

6th ISCA Speech Synthesis Workshop Scientific Committee Member, Bonn, Germany, August 22–24, 2007.

Interspeech 2006 satellite workshop “The Blizzard Challenge 2006: Evaluating corpus-based speech synthesis on common databases,” Co-organizer, Pittsburgh, PA, U.S.A., September 16, 2006.

Interspeech 2005 - Eurospeech “The Blizzard Challenge 2005: Evaluating corpus-based speech synthesis on common databases,” Special Session Co-organizer, Lisbon, Portugal, September 4–8, 2005.

5th ISCA Speech Synthesis Workshop Scientific Committee Member, Pittsburgh, PA, U.S.A., 14–16 June 2004.

IEEE 2002 Speech Coding Workshop Organizing Committee Member, Tsukuba, Ibaraki, Japan, 6–9 Oct, 2002.

Session Chair of numerous international conferences and symposiums, including ICASSP, Interspeech.

Reviewer of numerous journals and international conferences, including IEEE Transactions on Speech and Audio, Speech Communication, ICASSP, Interspeech.

PROFESSIONAL SOCIETIES

IEEE TASLP, Reviewer, 2023

ISCA Advisory Council, 2021-2024

NSERC, External Reviewer, 2020

ISCA Fellow Selection Committee, 2016.01.01-2016.12.31

European Science Foundation (ESF) Expert Reviewer 2016.10.20-2019.10.19

European Science Foundation (ESF) peer reviewer, 2007.05.01-2009.04.30

ISCA Fellow Selection Committee, 2015.01.01-2015.12.31

Scientific Advisory Board Member of an EPSRC Programme Grant “Natural Speech Technology (NST),” May 2011-April 2016.

IEEE Journal of Selected Topics in Signal Processing, Guest Editor, 2013.

IEEE Signal Processing Magazine, Guest Editor, 2012.

ISCA Advisory Council member, January 1, 2009–December 31, 2012.

Associate Editor of IEEE Transactions on Audio, Speech & Language Processing, Jan 1, 2009 – Jan 1 2012.

European Science Foundation (ESF) peer reviewer, 2007–present.

Ministry of Economy, Trade and Industry (METI), Japan, Project: “Development of technologies for the use of home information appliance sensors and human interface devices,” Academic Committee Member, 2007–2009

Acoustic Society of Japan, Member of The Board of Trustees, 2007–2008

Institute of Electronics, Information and Communication Engineers (IEICE), Guest Editor, 2005

The Robotics Society of Japan, Technical Research Groups for Robotic Audition, Committee Member, 2003–present

Acoustic Society of Japan, Associate Editor, 2003–2005

Institute of Electronics, Information and Communication Engineers (IEICE), Associate Editor, 2002–2006

Institute of Electrical and Electronics Engineers (IEEE), Signal Processing Society, Speech Technical Committee Member, 2000–2003

Institute of Electronics, Information and Communication Engineers (IEICE), Speech Technical Committee Member, 1999–2005

Japanese Society for Artificial Intelligence, Member of The Board of Trustees, 1998–2001

9 STUDENTS ADVISING

PHD STUDENTS AS PRINCIPAL SUPERVISOR

Heiga Zen, 2006
Ranniery deSilva Maia, 2006
Keiichiro Oura, 2010
Ryuta Terashima, 2010
Kei Hashimoto, 2011
Sayaka Shiota, 2012
Shinji Takaki, 2014
Akira Tamamori, 2014
Kazuhiro Nakamura, 2014
Kei Sawada, 2018
Takenori Yoshimura, 2018
Yukiya Hono ,2022

10 ANY OTHER RELEVANT INFORMATION

VISITING RESEARCHERS AND STUDENTS

Thomas Hain, University of Sheffield, U.K., 2018 (2 months)
Xin Wang, National Institute of Informatics, Japan, 2018 (3 weeks)
Simon King, University of Edinburgh, U.K., 2017 (2 months)
Amelia Gully, University of York, U.K., 2016 (2 months)
Anocha Rugchatjaroen, NECTEC, Thailand, 2016 (6 weeks)
Sittipong Saychum, NECTEC, Thailand, 2016 (6 weeks)
Rasmus Dall, University of Edinburgh, U.K., 2015 (6 months)
Hongwu Yang, Northwest Normal University, China, 2011 (1 year)
Ching-Tang Hsieh, Tamkang University, China, 2012 (3 months)
Alan W. Black, Carnegie Mellon University, U.S.A., 2012 (3 weeks)
Mikko Kurimo, Helsinki University of Technology, Finland, 2010 (2 months)
Shifeng Pan, Chinese Academy of Sciences, China, 2010 (3 months)
Ulpu Remes, Helsinki University of Technology, Finland, 2010 (8 months)

Linghui Chen, University of Science and Technology of China, China, 2010 (5 months)

Alan W. Black, Carnegie Mellon University, U.S.A., 2010 (2 weeks)

Jun Xu, Tsinghua University, China, 2008 (6 months)

Heng Lu, University of Science and Technology of China, China, 2008 (2 months)

Pascual Martínez Gomez, Universitat Politècnica de València, Spain, 2008 (1 year)

Alan W. Black, Carnegie Mellon University, U.S.A., 2008 (1 month)

Simon King, University of Edinburgh, U.K., 2007 (4 months)

Alan W. Black, Carnegie Mellon University, U.S.A., 2007 (2 weeks)

Alan W. Black, Carnegie Mellon University, U.S.A., 2006 (2 weeks)

Maria João Barros, Universidade Do Porto, Portugal, 2005 (3 months)

Christian Ants Weiss, Universität Bonn, Germany, 2005 (6 months)

Alan W. Black, Carnegie Mellon University, U.S.A., 2005 (1 week)

Fernando Gil Vianna Resende Junior, Universidade Federal do Rio de Janeiro, Brazil, 2004 (4 months)

Maria João Barros, Universidade Do Porto, Portugal, 2004 (1 month)

Fernando Gil Vianna Resende Junior, Universidade Federal do Rio de Janeiro, Brazil, 2003 (4 months)

Fernando Gil Vianna Resende Junior, Universidade Federal do Rio de Janeiro, Brazil, 2000 (2 months)

Junibakti Sanubari, Satya Wacana University, Indonesia, 1997 (1 year)