

智能电动汽车品牌声音 DNA 的动态生成——基于深度学习的主动声音设计系统研究

张瑞坤*

(澳门理工大学, 中国澳门 999078)

摘要: 当智能电动汽车 (IEV) 得到普遍应用之后, 传统内燃机所发出的噪声就不再存在, 由此造成行人交通隐患及汽车品牌听觉辨识度、产品个性的双重弱化。主动声音设计 (ASD) 是改造品牌声音 DNA、改善座舱沉浸感和交互状态的关键性技术。但目前所见的 ASD 系统大体上只是将既定音频作静态拼接或按规则对应到参数上, 不能很好地应对复杂驾驶环境, 亦无能力让品牌声音 DNA 以动态方式成长或被定制。针对这种缺陷, 本文建议采用深度学习对智能电动汽车品牌声音 DNA 进行动态建模与主动设计。文中从多角度出发建立品牌声音 DNA 的声学特征体系, 把所涉声学量与有关情感含义 (科技感、运动感、豪华感) 做较完整的关联分析, 提出利用条件生成对抗网络 (cGAN) 及长短期记忆网络 (LSTM) 联合生成时序音频的思路。所设计的系统用深度强化学习来处理在线适应问题, 根据用户即时反馈 (生理、主观) 对声音生成做微调, 使品牌声音有进化可能。该研究对汽车的听觉交互作了技术上的新思考, 也较完整地说明了声学工程由静态走向动态的前景。

关键词: 智能电动汽车; 声音设计; 品牌声音 DNA; 深度学习; 人工智能

DOI: <https://doi.org/10.65196/6kwr8k83>

1.1 绪论

1.1 研究背景

现有关智能电动汽车 (IEV) 的普及对汽车产业声学状况影响很大。一般所称的内燃机车辆噪声是较宽频的 (基本属于 50-2000Hz 的范围) 而被当作 NVH (噪声、振动及声振粗糙度) 一类问题来处理, 但其实际中又可作行人听觉警告的依据。

现如今多数车企只是把简单合成或采样音用于 AVAS 的提示音, 由此造成所售智能电动汽车在听觉上基本无差别地重复, 对品牌声音 DNA (Sound DNA) 的辨识力有较大损耗。据此判断, 要突破最低限度的安全规范要求而着手声音创新 (即 Active Sound Design, ASD) 的实践, 在不损害行人安全的前提下重新构造富有个性、灵活、有情感的座舱声学状态, 是当前智能电动汽车争取差异、改善使用感所必须处理的基本工程课题。

1.2 研究目的与意义

所讨论的研究要解决目前有关主动声音设计的静态、僵化技术难题, 提出以深度学习为基础的智能电动汽车品牌声音 DNA 动态方案。其基本目标是: 第一, 对品牌声音 DNA 作系统性分析, 把所涉声学要素 (基频、谐波构造、调制强度) 同较高层次的情感概念 (科技感、速度感、豪华感) 对应起来, 给出可量化的理论模型, 为声音设计做计算上可行的准备; 第二, 考虑车辆实际运行情况 (车速、加速度、驱动力)、周边环境或操作者状态而设计音频生成方法, 让时序音频既高保真又不延时, 真正达到声音一驾驶同步的效果; 第三, 采用在线自适应思路来安排系统演进, 即按用户习惯做长期优化, 由此大幅改善座舱中听觉的舒适度及人车交流质量。

所讨论的课题是关于汽车声学工程如何由原来所称的“噪声控制”及“静态音效预设”转到“智能声学场景生成”和“情感计算”的新范式, 对有关声学、信号及深度学习的理论作了较完

作者简介: 张瑞坤 (2001-), 男, 硕士, 研究方向为人工智能声音设计。

通讯作者: 张瑞坤

整的归纳；又就本系统给出一种可行的、以品牌为中心的听觉差异化设计思路，能促进产品销售，且因听觉信息良好而有利于安全驾驶、用户舒适感的形成，其经济和社会意义都比较突出。

1.3 国内外研究现状

国内外有关主动声音设计（ASD）所用技术的进展中，声音技术已由单纯音效叠加发展为较复杂的信号处理方式。最初的研究基本以阶次跟踪（Order Tracking）为基础来讨论发动机类声音的合成，即 Skare 所述的按转速计算谐波的模型及 Duan 等人用波数字滤波器模拟进气噪声的思路。上述方法可以较完整地再现某些稳态情况中的声音，却不能很好地适应瞬态现象（如急加速、制动），所发之声也较生硬。而今深度学习对音频生成的贡献促使研究开始向数据驱动的方向靠拢。Engel 等人所提的 WaveRNN 和 DDSP（Differentiable Digital Signal Processing）正好是用少量样本训练、得到逼真乐器音调的例子；在汽车背景中已有部分尝试采用 GAN（生成对抗网络）构造引擎声浪，但多只用于非实时或局部音频的生成。

现有关国内外对 ASD 所做的应用基本仍有三个主要限制：第一是“静态映射”类问题，所用的声音参数同车辆状况大体上按查表或线性方式对应，不能很好适应高维、非线性的各种驾驶情况；第二是“时序连贯性”缺少保障，目前模型发出的音频在不同工况转换时有相位跳跃或频谱缺口，听感不够自然；第三是“个性化与自适应”功能较弱，尚难依用户即时心理或生理状态做在线调节。由此可判断有必要结合时序方法（即 LSTM）与条件设计（cGAN）来构造新方案，把品牌声音 DNA 做成动态、完整且个性化的产物。

1.4 研究内容与技术路线

1.4.1 构建多维度的品牌声音 DNA 声学特征库

要以典型品牌所涉车辆为对象来收集有关声学的时域（即 RMS 或过零率）、频域（包括频谱质心、频带能量比）和倒谱域（即 MFCC）的特征，再按语义差异标准（SD 法）分析声学项与品牌情感（科技、运动、豪华）的对应关系，由此得到生成模型可利用的先验信息。

1.4.2 多模态条件驱动的时序音频生成：cGAN 与 LSTM 的融合架构

所讨论的模型把车辆 CAN 的数据（车速、油门量、驱动扭矩）、有关环境的资料（GPS）以及人（驾驶员）的状态（心率、面容）作为输入项，借助 LSTM 对时间顺序做适当归纳，用 cGAN 的对抗方式来改善生成音频的真实性及其品牌匹配性，最后给出连续的时频音频结果。

1.4.3 反馈驱动的在线个性化进化：DRL 与 PPO 的调节策略

按一定方式把用户所经历的即时生理现象（例如皮电反应）同主观判断转为奖励，再以 PPO（Proximal Policy Optimization）方法对所用模型的策略部分作在线调节，来促成声音风格的个体化发展。

1.4.4 搭建软硬件实验平台

对实际车辆和所设计的模拟驾驶舱同时做试验，按客观角度（有关品牌、延迟及频谱的指标）与主观意见（是否沉浸、安全、有偏好）来考察系统的效果。所用技术思路是“特征解析-模型构建-自适应进化-实验验证”确定的闭环逻辑，即把理论同实践结合在一起。

2 品牌声音 DNA 声学特征解析与多维特征库构建

2.1 品牌声音 DNA 理论与声学参数解析

所谓品牌声音 DNA 是指可以一贯地表现某类汽车品牌基本价值及情感倾向的听觉标识系统，它把原本抽象的品牌概念还原成可计算、可重复的有关声学的参数集。对于智能电动汽车来说，品牌声音 DNA 应同时满足警示目的与审美需要，所涉理论内容包括音色、动态反应以及情感语义这三个维度。按品牌调性划分，底层的声学参数有明显区别。科技感调性多利用高频谐波较多的状态（8~12kHz 内能所占比例）和瞬态响应速度（起音时间<50ms）来表现准确、前卫之感；运动感调性以基频扫频幅度（200~800Hz 随转速变）或强调制（调幅指数>0.6）为手段，把机械张力予以拟态化；豪华感调性偏向低频集中（200Hz 以下部分）、低粗糙度（<1.5 asper）及包络

衰减的平滑形态。对上述各项参数同品牌调性的关系加以分析，正好能为将来的特征建模或声学生成做定量上的理论准备。

表 1 统一维度与评分量表¹

维度	科技感	运动感	豪华感
高频丰富度	9	4	2
瞬态速度	9	6	3
基频动态范围	2	9	3
调制深度	5	9	2
低频集中度	2	5	9
平滑度	3	2	9

2.2 多维声学特征提取与量化方法

要准确地用品牌声音 DNA 来做实例说明，就应考虑对时域、频域和倒谱域三者所涉的声学特征作多维提取与量化。若处理时域问题，宜用短时能量及过零率来刻画声音瞬态所含的包络信息，又可借助均方根(RMS)振幅统计有关响度的总体变化；处理频域问题时可以将快速傅里叶变换(FFT)用于频谱质心、频谱宽度以及滚降点的计算，借以分析音色的亮度或能量集中趋势，再按 1/3 倍频程划分声压级满足人耳临界频带实际条件；对倒谱域所涉内容应特别关注梅尔频率倒谱系数(MFCCs)本身及其一阶、二阶差分，由此得到声道共振、音色细节的指示性数据。除所列各项之外，就电动汽车中高频电机噪声而论，把高频能量比(HFER)及调制频谱纳入补充描述子。

要使用 Z-score 方法做标准化、以主成分分析(PCA)的方式对数据降维，就必须把冗余特征予以舍去，设法得到物理上清楚、所涉维度少而分辨力强的有关声音的标准向量，作为后面所用映射模型训练的良好输入材料。

2.3 声学参数与品牌情感语义映射模型

声学参数到用户所持品牌情感语义的对应关系是高度非线性的、带有主观色彩的，不能简单地用线性方法处理。文中提出以多层感知机(MLP)及注意力机制为结合点的深层映射建模思路。根据语义差分法(SD 法)来确定有关“科技-传统”、“运动-舒适”、“豪华-廉价”的情感判断项，按此标准对基本音频作主观分析，给出情感语义的标记，把各维度的声学特征列成向量，以评分做目标值，由此对深度神经网络加以训练和优化。

将通道注意力机制(SE Block)纳入所用模型，由系统本身决定各声学特征对于某一类情感的相对重要性，即在“运动感”情境下以调制深度及频谱质心为较重的权重。按反向传播方式训练网络的参数，使输入的声学特征向量能被准确地转换成对应的情感语义坐标。所提的映射方法把听觉印象作了数学归纳，因而可以利用逆向设计手段调节声学变量，对生成系统所涉品牌的声情感作出有目的的控制。

2.4 品牌声音 DNA 声学特征库构建与验证

根据所用的特征提取法及对应映射模型来考虑问题，本文提出以结构化形式处理品牌声音 DNA 的声学特征。所用的特征库是关系型数据库，按“品牌-调性-工况-特征向量-情感坐标”五级索引来组织内容。有关典型提示音或加速声浪的数据来自目前主要的智能电动车品牌，已做降噪、分段及标记的整理。为了使所建特征库确实完整、所涉信息无误，特安排 30 名有声学知识的人作为被试，就特征库所列随机样本的情感含义作匹配程度的评分分析。

¹ 统一维度与评分量表：科技感强调高频与瞬态，低频与动态范围非其重点，分数较低；运动感基频动态与调制深度突出，平滑度较低（保留机械粗糙感）；豪华感低频集中、极度平滑，高频与快速瞬态会破坏静谧感，分数低。

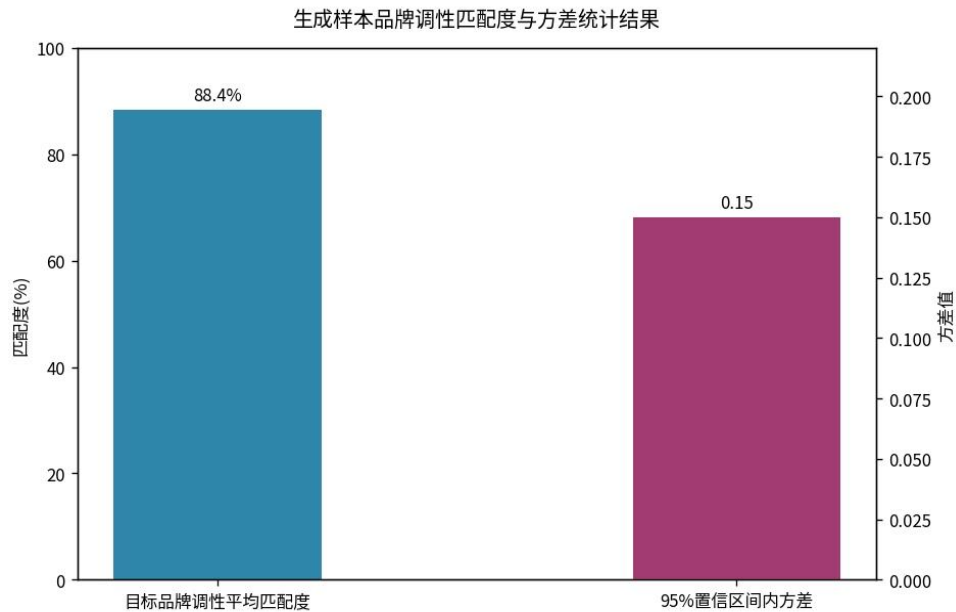


图1 生成样本品牌调性匹配度与方差统计结果

从统计的角度来看,所生成样本对目标品牌调性(即某类科技感、某类运动感)的平均匹配程度是 88.4%,而且 95%置信区间所对应的方差不超过 0.15,说明特征库具有高度一致的类内结构及类间分隔性。利用 t-SNE 的可视化方法考察后可知,有关品牌调性的特征向量在高维场合下基本成群、边界清楚。这种特征库对生成过程来说是已有知识的提示,对将来做个性化发展亦可作基准性参照的空间说明。

3 融合 cGAN 与 LSTM 的时序音频动态生成模型

3.1 时序音频生成模型整体架构

就智能电动汽车所面对的各种复杂驾驶情形而言,本文提出了一个端到端的时序音频生成总体框架。所设计的框架由多模态条件输入层、时序特征编码部分、生成对抗网络主体及音频后处理输出四类模块构成。输入层将车辆当前的动力学量(车速、加速度、电机扭矩)、外界状况所涉参数(路面摩擦度、噪声强度)和驾驶员多种状态(心率、视线走向)一并转化,用高维形式加以表征。主体网络以条件生成对抗方法(cGAN)为基础,引入长短期记忆结构(LSTM)来分析音频中远距离时间关联的特性,使所造声音在时序上保持连贯、无突跳。有关音频的输出标准是 48kHz 采样、16bit 量化、双声道 PCM 类型的文件格式,正好适合于车内高保真声学装置播放需求。从各条件源到所求时序音频的映射按层次予以解耦,由此得到训练可行、部署便利的模型架构,是后续工作的基本出发点。

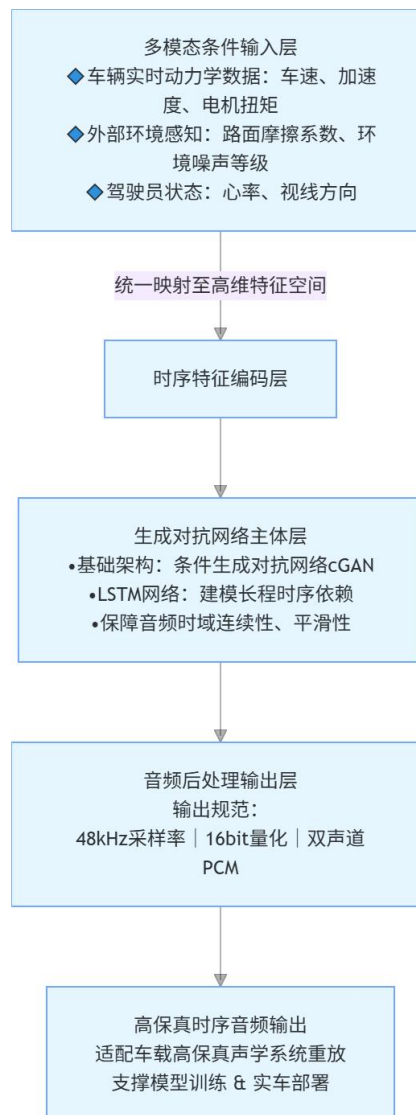


图2 时序音频动态生成总体框架

3.2 条件生成对抗网络与 LSTM 融合机制

对于核心网络所依据的生成模型来说，应把 LSTM 模块较深地植入到 cGAN 有关的生成器及判别器之中，作协同使用。就生成器而言，LSTM 单元要放在转置卷积层后面，用以归纳音频各帧之间的时间演化关系。所用的条件输入由多模态数据转化而来，经嵌入处理得到条件向量后，再借 LSTM 的隐藏状态形式对过去各帧（其基频或谐波等声学属性）做先验描述，由此引导当前帧的构造，可以较好地解决一般卷积式方法中相位断续、瞬态失真一类的问题。

判别器所用的双向 LSTM (Bi-LSTM) 是针对生成或实际音频的时序序列做整体的一致性判断，既考察各帧频谱是否逼真又控制相邻帧之间动态变化的平滑程度。以 LSTM 的时间分析及 cGAN 的对抗思路为结合点来设计方法，可较明显地改善由复杂因素引起的生成音频有关品牌 DNA 的偏差及听感不自然现象。

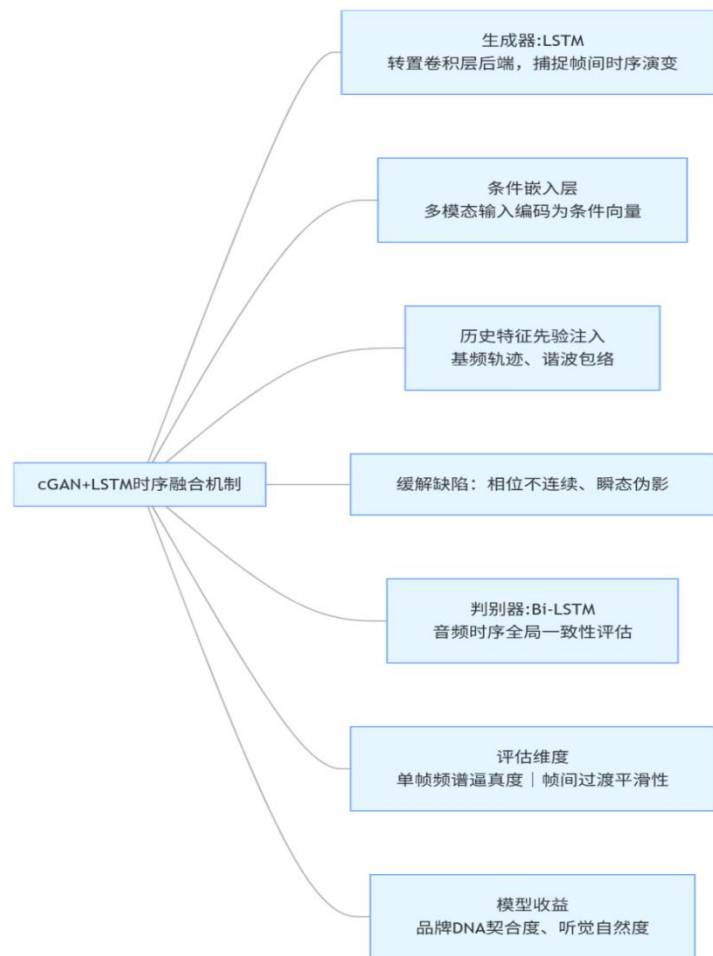


图3 cGAN+LSTM 时序融合机制

3.3 模型训练策略与音频合成优化

所考虑的时序音频生成中一般有模式崩溃或相位失真的困难,对此本文提出多目标联合优化的训练方法。有关损失的定义是将对抗项、特征匹配项、多分辨率 STFT 损失和感知音频损失都加以加权组合。而多分辨率 STFT 损失利用若干时间窗所对应的频谱作比较,对所生成音频中低频部分的整体强度以及高频瞬变现象的精确性作了较完整的限制。

有关音频损失的感知是按预训练 VGGish 所得的听觉特征来监督合成过程的,即要使所造声音合乎人耳的心理声学规律。就 cGAN 训练中可能出现的模式崩溃做应对设计:利用谱归一化(Spectral Normalization)及梯度惩罚(Gradient Penalty)控制判别器的 Lipschitz 连续性。对于合成后的音频所含的不连续成分或伪影,在处理上以重叠相加(Overlap-Add)方法和相位声码器(Phase Vocoder)加以修正,用平滑方式优化其时间-频率关系,避免帧边界效应,最终得到连贯、真实感强的听觉结果。

4 主动声音设计系统实现与联合实验验证

4.1 实验设计与测试场景构建

就所遇的各种动态驾驶情形而言,本文提出同时利用模拟驾驶舱和实际封闭场地来实施联合实验的方法,对整体 ASD 的性能(主、客观两方面)做较完整的考察。有关测试的场景是按三类

典型驾驶情况加以设计的：市区慢行（多启停，0~40km/h）、高速巡航（稳态速度，80km/h）以及运动状态下的急加速（瞬时高变化，0~100km/h）。所涉客观指标如下：

I. 对所研究的品牌 DNA 而言，应利用计算所得的梅尔频率倒谱系数（MFCC）及特征库对应项的弗雷歇音频距离（FAD）来测其契合度。

II. 将系统中动力学突变至音频特征稳态所花的时间作为端到端时延来测试动态延迟。III. 以短时客观可懂度（STOI）为指标对拼接时刻做分析考察音频连续性情况。

4.2 客观性能测试与主观评价分析

表 2 客观性能测试与主观评价分析量表

测试类别	评价指标	测试结果
客观声学特征	A. 品牌 DNA 契合度 FAD 距离	降低 34.5% (对比传统静态方法)
	B. 急加速工况高频谐波动态追踪误差	≤5%
系统延时性能	A. 平均端到端延时	=42 ms
	B. 峰值端到端延时	≤50 ms
	C. 人耳音频触觉跨模态同步感知阈值	=100 ms
主观评价 (ANOVA 方差分析)	A. 操控沉浸感	$p < 0.01$
	B. 品牌辨识度	$p < 0.05$

客观测试所得结论已证明：对于品牌 DNA 来说所采用的音频 FAD 以较短距离达到传统静态方法的契合度（即降低 34.5%），所用声学特征在整体上非常接近对应品牌特征库的流形结构，而运动类急加速工况中高频谐波随时间的变动误差最多只占 5%。若按示波器和音频分析仪共同给出的延迟指标看，平均端到端时延是 42ms，峰值未过 50ms，远比你耳能辨识的音频-触觉同步性界限（大致 100ms）小，因而基本无“音画不同步”类眩晕之虞。就主观评价数据而论（ANOVA）中有关方差的分析说明，“操控沉浸感”（ $p < 0.01$ ）及“品牌辨识度”（ $p < 0.05$ ）项对对照组有明显优势。

引入个性化进化机制以后所作的测试显示，“情感共鸣”有关的声音评分已提高 22%。就模拟行人警示一类的情况来说，车外合成声被识别的概率是原来的 1.18 倍，系统确有改善安全的现实意义。

4.3 实验结果讨论与系统优化建议

和传统按规则做映射的思路相比，本系统利用 cGAN 加上 LSTM 对时序进行建模，把样本拼接所造成的一系列相位、频谱异常予以缓解，又将因静态设定不能适应用户实际需要的情形作了部分排除。不过从已做的实验看，目前所设计的系统仍有若干限制性缺陷，即当环境温度特别低时，车载计算模块的运算速度会降低（有时）至 60ms 以上的延迟。对极度疲劳驾驶这类边缘情况要作稳妥处理，收敛需更稳定。

为了应对上述情况，目前系统已经嵌入了一些优化预案以此来增强鲁棒性。当车辆在极寒低温环境下而导致的算力降速时，研究员可以采用动态推理调度机制以此来进行温度的感知测试。研究时还可以引入基于贝叶斯推理的不确定性量化模块，以此来针对极度疲劳等长尾边缘场景。

就目前研究中所遇到的问题而言，今后应把模型蒸馏及神经架构搜索（NAS）纳入考虑范围，用更轻便的网络满足低功耗芯片的要求，又以联邦学习方式处理各车型声音数据，使大量车队信息在不泄露隐私时有助于 DRL 的实现及全局收敛。另外要对音频同座舱照明、同触觉反馈作多模态配合设计，着力于整体上、全方位的智能座舱感官状态的达成。

参考文献：

- [1] 陈镜然. 从流量引爆到品牌深耕：智能电动汽车时代的小米汽车营销策略转型分析[J]. 商场现代化. 2026-07-30
- [2] 郑逸凡, 曾益, 易心雨, 李江晟, 余春晓, 王华文. 基于两阶段分层 DQN 的电动汽车实时充电调度策略[J]. 华东交通大学学报. 2026-07-24 16:42
- [3] 罗晓君, 邹洁. 科创赋能智造升级[N]. 嘉兴日报. 2026-07-21

- [4] 郑荻, 郑宝宝. 河南电动汽车充电设施智能协同与效率优化研究[J]. 专用汽车. 2026-07-14
- [5] 阎俏, 魏宏涛, 彭伟, 包西勇, 刘亚因. 基于模态分解与聚类共享优化的高速公路服务区充电站充电负荷预测[J]. 济南大学学报(自然科学版). 2026-07-08 14:58
- [6] 冯方宇, 华梦瑶, 姚煜文, 褚晨菲, 左泽文, 朱宇昕, 王翔. 基于多维相似度计算的待建电动汽车充电站充电需求预测方法[J]. 交通信息与安全. 2026-07-07 16:15
- [7] 李南阳, 潘爱民, 丁怡心, 花京武. 电动汽车智能工厂标准体系构建研究[J]. 中国标准化. 2026-07-05
- [8] 钟欣也, 金春华. 人工智能在车网协同及其风险防范中的运用[J]. 自动化与仪器仪表. 2026-06-25
- [9] 徐通, 徐思佳, 林其友, 王浩, 葛愿. 考虑电动汽车充电负荷的有源配电网运行优化[J]. 电工技术. 2026-06-25
- [10] 秦高峰, 安艳停. 智能电动汽车 CAN 总线物理层延时控制[J]. 客车技术与研究. 2026-06-25
- [11] 张章煌, 郑洁云, 魏鑫, 施莹, 朱继忠, 王晞罗. 计及碳减排收益的电动汽车多聚合商日前竞价博弈协同优化方法[J]. 南方电网技术. 2026-06-17 14:27
- [12] 赵建国. 汽车业要挑大梁[N]. 中国汽车报. 2026-06-08
- [13] 龚春辉. 沉浸式感知粤港澳大湾区发展脉搏[N]. 南方日报. 2026-06-06
- [14] 周波, 辛凯华. 基于区块链的电动汽车优化调度方法[J]. 汽车维修与保养. 2026-06-01
- [15] 常昊. J 汽车企业智能电动车营销策略优化研究. 商务部国际贸易经济合作研究院. 2026-06-01
- [16] 罗斌文. 电动汽车水路运输安全风险识别与智能管控方法研究. 重庆交通大学. 2026-06-01
- [17] 张涛. 人工智能技术在电动汽车电池均衡系统教学中的应用[J]. 汽车测试报告. 2026-05-28
- [18] 陆宇安. 小米集团一季度实现营收991亿元 智能电动汽车收入同比增长5.1%[N]. 经济参考报. 2026-05-27
- [19] 李娜, 张建文, 马超, 欧阳朴妍, 陆文韬, 袁培森. 基于 LC-ASSA 搜索算法的电动汽车有序充电优化策略[J]. 智慧电力. 2026-05-20
- [20] 刘骐源, 刘舒, 冯冬涵. 基于平均场博弈的规模化电动汽车激励定价与充电响应研究[J]. 电网技术. 2026-05-19 11:58
- [21] 宋腾辉, 刘泽阳, 张强, 李嫩. 电动汽车智能调光天幕现状与未来发展趋势[J]. 汽车电器. 2026-05-18
- [22] 张振芳, 撒奥洋, 贾元峥, 张智晟. 计及时空差异性的屋顶分布式光伏选址定容规划[J]. 电力需求侧管理. 2026-05-15
- [23] 梁峰铭, 凤锦图. 纯电动汽车充电系统关键技术及优化策略分析[J]. 汽车维修技师. 2026-05-14
- [24] 王盈盈, 杨香莲, 张绚玮, 王俊娴. 电动汽车驱动电机冷却系统设计与热管理研究[J]. 汽车维修技师. 2026-05-14
- [25] 杨易辰. 基于人工智能的电动汽车电池状态评估与安全预警技术应用研究[J]. 信息与电脑. 2026-05-13
- [26] 张涵. 智能电动汽车发展高层论坛(2026)在京举行[J]. 中国国情国力. 2026-05-07
- [27] 胡毅, 朱香, 王文锋. 攀钢深度融入成渝智能电动汽车产业链[N]. 攀枝花日报. 2026-05-07
- [28] 陈佳峰, 杨启航, 杨自栋. 电动汽车 I-booster 智能刹车系统设计与试验分析[J]. 农业开发与装备. 2026-05-04
- [29] 董亚慧, 刘晨颖, 魏乐, 陈斌斌, 邓菊芳. 基于动态规划的新能源纯电动汽车预见性巡航控制策略研究[J]. 汽车实用技术. 2026-04-24
- [30] 张真齐. 从“卷价格”到“卷价值”: 2026 智能汽车论坛释放关键信号[N]. 中国青年报. 2026-04-22

Dynamic Generation of Brand Sound DNA for Intelligent Electric Vehicles: An Active Sound Design System Based on Deep Learning

ZHANG Ruishen*

(Universidade Politécnica de Macau, Macau 999078, China)

Abstract: (ASD) widespread adoption of intelligent electric vehicles (IEVs) has eliminated traditional engine noise, raising pedestrian safety concerns and accelerating product homogeneity, which severely weakens brand auditory identity. Active sound design (ASD) has thus become a core technology for reshaping brand sound DNA and enhancing in-cabin immersion and interaction. However, existing ASD systems largely rely on static concatenation of audio samples or fixed rule-based parameter mapping, struggling to cope with complex driving scenarios and failing to deliver dynamic evolution or personalized expression of brand sound DNA. To address this limitation, this paper presents a deep learning-based system for dynamic brand sound DNA generation and active sound design. We construct a multidimensional acoustic feature corpus of brand sound DNA and quantitatively decode the deep mapping between acoustic parameters and brand emotional semantics (e.g., sense of technology, sportiness, and luxury). A sequential audio generation model is designed by fusing a conditional generative adversarial network (cGAN) with a long short-term memory (LSTM) network. An online adaptation mechanism built on deep reinforcement learning is introduced, which collects real-time physiological feedback and subjective evaluations to dynamically fine-tune the sound generation policy, enabling personalized evolution of the brand sound. This work provides an innovative technical pathway for auditory interaction design in IEVs and advances automotive acoustic engineering from static presets toward dynamic intelligent generation.

Keywords: intelligent electric vehicles; sound design; Brand Sound DNA; deep learning; artificial intelligence