

## A Three-Dimension Quantitative Model about Tone Sandhi

---

### Abstract

Paper attempts to design an model to interactively cooperate left-dominant and right-dominant tone sandhi patterns.

The correlation between Wu Sino-Japanese and Amoy affiliated nasals initials is presented by Hu.F(2005). Based on Hu's experimental phonetic investigation, this paper proposes a three-approaches multimodal in computational modal design, with acquisition background.

Firstly, synthesized child acquisition data of tone sandhi, a multimodal construction is renewed to solve complex real-word data transformation problem by graduate steps. And, deep neural network is applied in describing tone circle phenomenon of Min tone sandhi.

Simultaneously, a more complicated apparatus is depicted to verify correlation between ancient Chinese ingredients and Japanese geminates acquisition accuracy (perception as acquisition refer indicator from Ren.H; Mariko.K 2018). With machine learning data of Sino-Japanese and Southern Min (Wang.C;Liang.H;He.M 2024), interactive combination between machine learning and traditional linguistics archaeological chronology is described as a well-rounded model filling up blank board of language field.

In nutshell, deeply discussing reconstruction towards cascading stages amid Amoy pre-nasals [<sup>m</sup>b,<sup>n</sup>d, <sup>ŋ</sup>g] evolutions: m<sup>b</sup>-<sup>m</sup>b- b (Hu.F 2005), around mathematical nature of Lagrangian k/d operators' transformation in Maximum Entropy Model(i.e use Gradual Learning Algorithm by Zhang.J&Lai.Y 2008) and dialectical distance quantitated in Mutual Entropy Model(Wang.B 2020), the phonotactic productivity of affiliated nasals and codas is multimodally ranked with quantitative values, using filtering workflow of optimality.

Future-visionally, linking with almost 2136 acquisition tokens(12 testers × 89 syllables) and 22 computational items of comparative data from Mutual Entropy Model in machine learning (Wang.C;Ge.C;Peggy.M 2024; Wang.C et.al 2024),a three-lines experimental design combining Geographic Typology Model (Hashimoto 1978) is presented to discuss the future direction of multimodal database, with help of LLM.

## LITERATURE REVIEW

## Nasal-affiliation:

胡方（2005）认为厦门话[b g l]声母的声学特性，指出它们实际上是鼻冠音[mb ŋg nd]，而不是传统认为的一般浊塞音和边音，进一步探讨了这些声母的历史成因，认为此类声母都是鼻音塞化的结果（闽南方言、粤语、客家话、晋语等），对此音变的机理进行了语音学上的解释（鼻冠音、后塞鼻音、鼻音塞化的对比区别）

系统地分析厦门话 b l g 声母的声学特性,并在此基础上讨论一些相关的语音、历史演变等

结论：

历史演变：厦门话鼻冠声母 (ᵐb ᵑd ŋg) 来自古鼻音声母 (\*m \*n \*ŋ)

共时分析：m-b 中间有两个阶段，mᵇ-ᵐb；潮汕话在第一阶段，厦门话在第二阶段

## 浊音声母

b-古次浊声明母 m

g-微母和疑母 ŋ

l-音色接近 d，介乎「l」和「d」之间-古次浊声泥、来、日母(娘)

bl (d) g-在口元音前

mnŋ-鼻化元音前 (~)

互补对位，没有音位对立关系

位置区别：音节间的鼻冠音往往接近一般浊塞音 (b l g)；普通鼻音还是保留鼻音 (m n ŋ) 面 mi 食

厦门话鼻音声母字在日译吴音和汉音中有不同读音的现象.不同材料中的鼻音的这个相同的特征有共同的来源呢还是各自独立发展的结果?

（严棉；综述）

古音小镜：

这个问题，厦门话鼻冠音产生的机制来看，答案偏向于鼻音带塞音特征 (mᵇ)是后起的。是各个方言自身独立发展的结果的可能性比较大。

## 方法 1

用 affiliate 运用于 coda

Affiliate: 量化定性 (distance, 演化速度, f1/f2 鼻化元音配位-网格)

Coda

计算入声韵尾与相同语言点的 distance 和演化速度, 与相对范畴进行对比, 是否能通过量化模型比对进行定性分析

代入习得 coda 正确率进行判断辅助定量模型校准

OT 框架和算法描写

入声韵尾

-t -k -p

《北江盆地——宗族聚落的形态与发生史研究》2012: 本书有关宗族、聚落的史地研究专著, 以浙东北江盆地的研究为模型, 探索“近世型宗族”的源流与发展脉络, 重建了“近世型宗族”的发生发育历程。该研究为浙东地区宗族聚落演化史研究提供明确、细致的时间标尺, 对研究今日宗族与聚落面貌及其长期演化趋势提供必要的理论参考与个案素材。

本研究基于丁邦新 (2002) 提出的“吴闽关系说”, 以及日译吴音来源出发, 以汉语中古音为参照, 从历史音韵学角度梳理各自的历时层次, 并分别选取合适的语料范围, 利用“互信息熵”模型对中古汉语、方言及域外方言进行存古成分的音韵量化, 通过数据分析、类型地理学模型的构拟等考察汉字音吴音与闽南语方言之间的相似度。同时, 基于互信息熵的存古程度计算, 经过数据处理后的语言点比对相似度  $S$  又能转化为语言年代学下的语源词留存率  $C$  (proportion of retained cognate) 进而比出地域空间下不同语言在历史特定谱系点的演变速度  $R$  (constant rate of retained cognate)。

本研究是基于机器学习视角对互信息熵模型在音系音韵研究的应用, 即选取适当的量化模型对声母/韵母‘存古度’的一种定量表述, 研究中关于衡量汉语和域外方言层次距离的影响因子通过语料库的量化计算和不同模型因子的联立转换获得了准确率提升, 使得方言存古度的信息熵模型和年代语言学视角下的统计理论得以获得数据上的交互结合。

本研究首先基于胡方（2005）对厦门话鼻音  $m^b$ - $n^b$ - $b$  三阶段发展讨论，即厦门话鼻冠音处于鼻后塞音逆同化第二阶段-第一阶段的延伸-该一结论，对比另一中古汉语的镜像转写方言-日语汉字音，结合日语波行声母万叶假名早期读音 P 音、以及  $P \rightarrow F \rightarrow H$  的历史演变考证（曾婧 2012），以桥本万太郎（1987）的语言地理类型学视角聚焦闽南语和汉字音这一域外方言的形态发展框架。同时兼之与机器学习、年代语言学（chronological linguistics）结合的“方言量化存古”模型辅助。

厦门话鼻冠音的位置作用（syllabic non-syllabic）对其浊音化，即向第三阶段转变，存在显著影响。由于鼻冠音声母、入声韵尾在语音分析软件的数据采集精度不如元音，本研究基于前人结论，通过机器学习、互信息模型的存古度量化整合方法，为鼻冠音声母、入声韵尾历史演变提供进一步数据（语言距离度 D）支持。

最后，联系年代语言学的词源考古计算（R FORMULAR），多模型的量化构建中，习得数据也将成为讨论的一环（王纯等 2024；任宏昊 2018）。

针对鼻冠音浊化的三重阶段发展，以及日语汉字音、中古汉语、闽南语在万叶假名  $P \rightarrow F \rightarrow H$  读音演变中的综述讨论，方言特殊声母、韵尾的语音形态演变框架，将通过数量计算模型因子（S R）以综合优化流程选出。

日语波行声母在万叶假名早期读 p 音的闽南语读音证据（p 音）