

Galaxy and Mass Assembly (GAMA): probing galaxy-group correlations in redshift space with the halo streaming model

Qianjun Hang,^{1,2★} John A. Peacock¹, Shadab Alam,¹ Yan-Chuan Cai,¹ Katarina Kraljic^{1,3},
Marcel van Daalen⁴, M. Bilicki⁵, B. W. Holwerda⁶ and J. Loveday⁷

¹*Institute for Astronomy, University of Edinburgh, Royal Observatory Edinburgh, Blackford Hill, Edinburgh EH9 3HJ, UK*

²*Department of Physics and Astronomy, University College London, Gower Street, London WC1E 6BT, UK*

³*CNRS, CNES, UMR 7326, Laboratoire d'Astrophysique de Marseille Aix Marseille Université, Marseille, 13388 CEDEX 13, France*

⁴*Leiden Observatory, Leiden University, PO Box 9513, NL-2300 RA Leiden, the Netherlands*

⁵*Center for Theoretical Physics, Polish Academy of Sciences, al. Lotników 32/46, P-02-668 Warsaw, Poland*

⁶*Department of Physics and Astronomy, University of Louisville, Natural Science Building 102, Louisville, KY 40292, USA*

⁷*Astronomy Centre, University of Sussex, Falmer, Brighton BN1 9QH, UK*

Accepted 2022 September 7. Received 2022 August 24; in original form 2022 June 10

ABSTRACT

We have studied the galaxy-group cross-correlations in redshift space for the Galaxy And Mass Assembly (GAMA) Survey. We use a set of mock GAMA galaxy and group catalogues to develop and test a novel ‘halo streaming’ model for redshift-space distortions. This treats 2-halo correlations via the streaming model, plus an empirical 1-halo term derived from the mocks, allowing accurate modelling into the non-linear regime. In order to probe the robustness of the growth rate inferred from redshift-space distortions, we divide galaxies by colour, and divide groups according to their total stellar mass, calibrated to total mass via gravitational lensing. We fit our model to correlation data, to obtain estimates of the perturbation growth rate, $f\sigma_8$, validating parameter errors via the dispersion between different mock realizations. In both mocks and real data, we demonstrate that the results are closely consistent between different subsets of the group and galaxy populations, considering the use of correlation data down to some minimum projected radius, $r_{\min}$. For the mock data, we can use the halo streaming model to below $r_{\min} = 5 h^{-1}$ Mpc, finding that all subsets yield growth rates within about 3 per cent of each other, and consistent with the true value. For the actual GAMA data, the results are limited by cosmic variance: $f\sigma_8 = 0.29 \pm 0.10$ at an effective redshift of 0.20; but there is every reason to expect that this method will yield precise constraints from larger data sets of the same type, such as the Dark Energy Spectroscopic Instrument (DESI) bright galaxy survey.

Key words: gravitation – galaxies: groups: general – large-scale structure of Universe.

1 INTRODUCTION

The large-scale structure in the galaxy distribution has a long history of providing cosmological information. The first constituents of the inhomogeneous galaxy density field to be identified were the rich clusters, which today we see as marking the sites of exceptionally massive haloes of dark matter. Proceeding down the halo mass spectrum, we find progressively less rich groups of galaxies, leading to systems dominated by a single L_* galaxy, such as the Local Group (e.g. Wechsler & Tinker 2018). All these systems have been familiar constituents of the Universe since the first telescopic explorations of the sky, but it took rather longer to appreciate that they were connected as part of the cosmic web of voids & filaments (see e.g. Peacock 2016, for some selective history). In part, the history here showed a complex interaction of theory and observation, since redshift surveys through the 1980s lacked the depth and sampling to reveal the cosmic web with complete clarity. For a period, it was therefore a question of asking whether the real Universe displayed

the same structures that were predicted in numerical simulations of structure formation in the Cold Dark Matter model (Bond, Kofman & Pogosyan 1996). But since those times, there has been an increasing confidence that galaxy groups are indeed particularly extreme non-linear points in the general field of cosmic density fluctuations, and this makes them interesting in two ways. First of all, groups are readily identified in galaxy surveys, providing a relatively robust data set (Eke et al. 2004; Robotham et al. 2011). Secondly, their non-linear nature makes them an informative probe of theory. Modelling non-linear behaviour is by its nature challenging compared to linear theory, but by studying structure formation further into the non-linear regime, we have the chance to test the robustness of our cosmological conclusions.

Our specific aim in this direction is to use galaxy groups as a probe of the cosmological peculiar velocity field. Such deviations from uniform expansion must exist through continuity, and density concentrations such as groups should be associated with an average infall velocity in regions surrounding the groups. The amplitude of these velocities depends in part on the strength of gravity on cosmological scales, and the peculiar velocity field has thus increasingly been seen as a means of probing the nature of gravity

★ E-mail: e.hang@ucl.ac.uk

and testing alternative theories (e.g. Jain & Khoury 2010). Although, it is possible to probe peculiar velocities directly using absolute distance indicators (Davis et al. 2011), the most powerful tool has been Redshift Space Distortions (RSD). These arise inevitably in the study of the 3D galaxy distribution because the distances to galaxies observed on the sky are inferred from their redshifts, z , via the standard relation

$$d(z) = \int_0^z \frac{c \, dz'}{H(z')}, \quad (1)$$

where c is the speed of light and $H(z) = \dot{a}/a$ is the Hubble parameter. But this equation does not give the true distances, because Doppler shifts from the peculiar velocities modify the observed redshift: $1 + z \rightarrow (1 + z)(1 + v_r/c)$, where v_r is the radial component of the peculiar velocity. If we then use the observed redshift as if it were a true indicator of distance, we obtain a distribution of galaxies in ‘redshift space’ – in which the apparent properties of galaxy clustering are distorted in an anisotropic way.

These distortions have characteristics that depends on scale: outside large density concentrations, galaxies fall coherently together under gravity; while the orbital velocities inside dark matter haloes are effectively randomized. The latter effect convolves the redshift-space density field in the radial direction, leading to the characteristic radial elongations of high-density regions known as ‘Fingers of God’ (FoG). RSD due to coherent flows in the linear regime were first studied by Kaiser (1987). The growth factor f is defined by

$$f \equiv \frac{\partial \ln \delta}{\partial \ln a} \simeq \Omega_m(z)^{0.55}, \quad (2)$$

where δ is the matter overdensity, a is the expansion factor, and Ω_m is the matter fraction; the approximation for $f(\Omega_m)$ only applies for flat Λ CDM models in standard gravity (Lahav et al. 1991; Wang & Steinhardt 1998; Linder 2005). In Fourier space, and in the small-angle limit of a distant observer, the matter power spectra in redshift space and in real space are related by

$$P_m^s(k, \mu) = P_m^r(k)(1 + f\mu^2)^2, \quad (3)$$

where μ is the cosine of the angle between the wave-vector k and the line of sight. This simple equation was highly influential from its first appearance (Kaiser 1987), as it offered the chance of measuring Ω_m from measuring the RSD anisotropy. But eventually goals shifted as Ω_m became very well determined from other routes (especially the Cosmic Microwave Background). Following Guzzo et al. (2008), the modern view is therefore to emphasize that the growth rate for a given density is also proportional to the strength of gravity, so that RSD can be used as a test of theories of gravity.

The RSD signal has been measured by a number of surveys, including the 2dFGRS at $z \simeq 0.2$ (Peacock et al. 2001; Hawkins et al. 2003); the 6dFGS at $z \simeq 0.1$ (Beutler et al. 2012); the SDSS BOSS & eBOSS surveys at $z \simeq 0.6$ (Reid et al. 2012; Alam et al. 2017, 2021a); and at $z \simeq 1$ by the 8 m VVDS and VIPERS surveys (Guzzo et al. 2008; Pezzotta et al. 2017). For the GAMA survey at $z \simeq 0.4$, aspects of RSD were studied by Blake et al. (2013) and Loveday et al. (2018), who measured the pair-wise velocity dispersion to small scales and as a function of luminosity. The above studies all focused on galaxy autocorrelations.

The challenge in modelling RSD is that truly linear modes are rare. In observation, large scales are affected by cosmic variance due to the finite survey volume. McDonald & Seljak (2009) proposed the use of multiple tracers in order to overcome cosmic variance, although in practice the improvement is slight (Blake et al. 2013). To gain more information, one needs to probe smaller scales, where

the effect of non-linearity can systematically bias the results (de la Torre & Guzzo 2012).

One possible solution to this dilemma is to use galaxy groups to probe the velocity field. Due to the small random virial velocity of the central galaxy at the group centre, the coherent large-scale infall velocities of groups are dominant down to intermediate and small scales. The group autocorrelation would thus have reduced FoG, aiding the extraction of the linear growth rate (Padilla et al. 2001; Mohammad et al. 2016). In practice, the group catalogue in GAMA is sparse, with a number density of $4.3 \times 10^{-3} h^3 \text{ Mpc}^{-3}$ between $0.1 < z < 0.3$, and measurements of the autocorrelation will have high statistical noise. The cross-correlation between groups and galaxies is thus an intermediate route, which effectively improves the statistical power while still reducing the non-linear pairwise velocities at small scales. The clustering of GAMA groups has been recently studied in Riggs et al. (2021), and this work extends this study to further subsets of the data, concentrating in more detail on their different RSD signals.

Our aim here is thus to test the robustness of RSD methods down to small or intermediate scales using multiple tracers involving galaxy groups. By cross-correlating galaxies of different colours, and groups in different mass bins, we examine the consistency of the inferred cosmological results between the subsamples. In order to pursue this investigation, we develop a new model for RSD in cross-correlation, involving a combination of the halo model and the streaming model, which we implement by including some information taken from mock data. Throughout the analysis, we adopt the WMAP7 Cosmology (Komatsu et al. 2011) with $\sigma_8 = 0.81$, $\Omega_m = 0.27$, $h = 0.70$, and $n_s = 0.967$, consistent with the mock catalogue.

The GAMA data set and its mocks are detailed in Section 2 followed by Section 3.1 where we introduce the statistics for measuring the 2-point function in the data. In Section 3.2, we present the resulting 2D correlation function measurements for subsamples. In Section 4, we discuss the theoretical modelling of RSD in galaxy-group cross-correlations, and in Section 5 we confront this modelling with real and mock GAMA data. The models are validated in Section 5.1 via detailed comparison with the GAMA mocks, where we establish the scales to which the different theories can work without bias; we present the fitting of the real GAMA data in Section 5.2. Finally, we summarize the work in Section 6. We include our Appendices A – E in the online Supplementary materials.

2 GAMA DATA AND MOCKS

This analysis is based on the Galaxy And Mass Assembly (GAMA) spectroscopic survey. This was conducted using the 2dF facility at the Anglo-Australian 4-m telescope over 210 nights between 2008 and 2014, accumulating spectra of 265 958 distinct galaxies. Together with existing data, this yielded a catalogue of 330 542 redshifts over five survey fields totalling 250 deg^2 , with a mean redshift of $z \simeq 0.2$ (Driver et al. 2022). The three main fields near the equator, G09, G12, and G15 are used here, each covering an area of $12 \times 5 \text{ deg}^2$. The survey has an extinction-corrected r -band flux limit of $r < 19.8$, based on SDSS photometry.

The overall redshift completeness of the GAMA equatorial region is 98.5 per cent: this high completeness was achieved by a large number of repeated visits to 2dF fields covering the survey area in different ways. This property is greatly advantageous for small-scale galaxy and group studies compared to much larger surveys such as BOSS, where fibre collisions can lead to substantial undercounting of close galaxy pairs and thus bias the measured galaxy 2-point correlation function (Guo, Zehavi & Zheng 2012).

The present analysis uses the DR3 data release (Baldry et al. 2018), which differs slightly from the final data release, DR4 (Driver et al. 2022). DR4 implements revised flux completeness limits through the use of new KiDS photometry. The original *SDSS* limit of $r < 19.8$ was trimmed in DR4 to $r < 19.58$ for 98 per cent completeness in the equatorial fields. Galaxies for the present study are selected from the SpecObjv27 DMU (Data Management Unit), with CMB frame redshifts adopted from DistancesFramesv14. We apply the following criteria: redshift quality $nQ \geq 3$, angular completeness mask > 80 per cent, and visual classification $VIS_CLASS = 0, 1, 255$.¹ The spectroscopic redshifts are computed by the code *runz* (Driver et al. 2011), which has a 1σ redshift error of 50 km s^{-1} in terms of peculiar velocity (Liske et al. 2015). In order to compute correlation statistics, it is essential to accompany the galaxy sample with a knowledge of the survey selection in angle and redshift. As usual, this information is captured by a random catalogue *Randomsv02* of fictitious unclustered galaxies; this catalogue was generated by Farrow et al. (2015) from the actual GAMA galaxy catalogue using a modified method following Cole (2011). The idea of this method is to clone each galaxy n times and distribute them randomly within the maximum volume V_{max} that the galaxy can be observed given the survey magnitude limits

$$n = n_{\text{clones}} \frac{V_{\text{max}}}{V_{\text{max}, \text{dc}}}, \quad (4)$$

where $n_{\text{clones}} = 400$ is the total number of randoms divided by data, and $V_{\text{max}, \text{dc}}$ is the maximum volume weighted by overdensity $\Delta(z)$. This method is iterated until $\Delta(z)$ converges, and the redshift distribution of the resultant random catalogue is smooth without large-scale features (see fig. 4 in Farrow et al. 2015).

The official GAMA group catalogue (G3C) was constructed by Robotham et al. (2011). Most of the groups are found within $z \lesssim 0.35$ (see fig. 16 in Robotham et al. 2011): thus we impose a redshift cut $0.1 < z < 0.3$ for the groups. The group catalogue is derived using an anisotropic friends-of-friends (FoF) algorithm calibrated against an N -body mock catalogue. However, in order to have consistently defined groups in the GAMA mocks (see Section 2.3), we do not use the official G3C catalogue. Instead, we apply a similar FoF group finder algorithm due to Treyer et al. (2018) to both data and mocks (see Kraljic et al. 2018 for an application to GAMA). The main difference between the two algorithms is the parametrization of the linking length, and a detailed description of the algorithm and assessment of the group reconstruction quality can be found in the appendix of Treyer et al. (2018).

In addition to the above selections, we further split galaxies and groups into subsamples based on galaxy colour and group mass. The number of selected galaxies and groups in each GAMA field and for each subsample is summarized in Table 1. We describe the selection in more detail below.

2.1 GAMA galaxy colour selection

Galaxies are divided into two populations that are known to have distinct clustering properties: the ‘red’ galaxies, which tend to be older, with little or no active star formation, and the ‘blue’ cloud, where galaxies are younger with active star formation. To obtain the galaxy colours, we use the extinction corrected *SDSS* magnitudes

¹ $VIS_CLASS=0$: Not visually inspected but suspicious based on *SDSS* flags; $VIS_CLASS=1$: Visually inspected and a valid target; $VIS_CLASS=255$: Not visually inspected but should be OK based on *SDSS* flags.

Table 1. Number of selected galaxies and groups from GAMA fields with redshifts $0.1 < z < 0.3$ and flux limit $r < 19.8$. Galaxies are split into two colour classes, red and blue, by equation (5). Groups are split into three stellar mass bins: 40 per cent (LM), 50 per cent (MM), and 10 per cent (HM) by mass ranking from low to high, covering the mass range $\log_{10}(M_*/h^{-2} M_{\odot}) = 9.5\text{--}12.5$.

Number of		G09	G12	G15
Galaxies	Blue	17 335	18 719	19 053
	Red	20 584	22 155	21 141
	Total	37 919	40 874	40 194
Groups	LM	1877	2084	2054
	MM	2347	2606	2569
	HM	470	522	514
	Total	4694	5212	5137
G3C	Total	4937	5367	5358

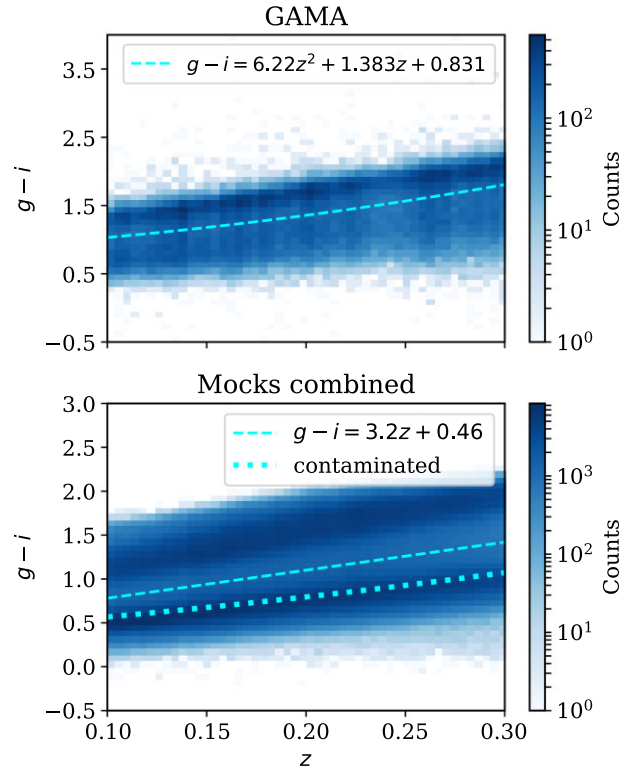


Figure 1. Distribution of the $g - i$ colour of galaxies in the redshift range $0.1 < z < 0.3$, for the real GAMA data (upper panel) and the average of 25 mocks combined (lower panel). Red and blue populations are separated by the dashed cyan lines. The cut in GAMA is chosen such that both GAMA and mocks have similar red and blue fractions at any given redshift. The dotted lines show the cut for an alternative ‘contaminated’ red sample (see text).

from the TilingCatv46 DMU. It is, however, non-trivial to separate the galaxy population into these subsets, because the colour distribution is continuous without gaps: elaborate approaches have been discussed in e.g. Taylor et al. (2015). For the purpose of this study, we adopt a simple quadratic cut in the apparent $g - i$ colour versus redshift plane:

$$g - i = 6.220z^2 + 1.383z + 0.831. \quad (5)$$

A cut of this form is motivated empirically by the apparent bimodality in the colour–redshift plane, as shown in Fig. 1. The precise location of the cut was adjusted in order to match the red and blue fraction

at each redshift in the GAMA data with the corresponding result in the mocks (which are discussed below in Section 2.3). The overall fraction of red or blue galaxies is very close to 0.5, and it changes only slightly with redshift: at the low redshift end, the red and blue fractions are similar, while towards higher redshifts, the fraction of red galaxies increases mildly until $z \sim 0.2$, and the difference in the red and blue fraction then becomes small at $z \sim 0.3$. We create random catalogues for the red and blue galaxy subsamples where required by applying this smoothly varying colour balance to the redshift distribution of the main random catalogue.

2.2 GAMA group mass selection

Groups are accepted with ≥ 2 group members, and the centre of the group is determined by the most (more) massive member in terms of stellar mass. The 2-member systems make up 66 per cent of the total groups in the GAMA data, but are likely to have poor fidelity. Thus, we emphasize that having the same group finder algorithm for the data and mocks is vital in order for these low-fidelity groups to be comparable. There are several approaches for determining the group centre. The simplest choice is to select the most massive member to be the central galaxy, and assume that it overlaps with the halo centre. Other approaches include determining a weighted centre by averaging over the positions of the group members, or iteratively excluding members that are most distantly separated (see e.g. Robotham et al. 2011). The iterative centres are used in the G3C catalogue, and it is shown in Robotham et al. (2011) that the agreement with using the brightest group galaxy (BCG) as group centre is 95 per cent for groups with $N \geq 5$, and that both BCG and iterative centres give highly consistent results for $2 \leq N \leq 4$ compared with the mock, and that the BCG centres are only degraded by about 3 per cent compared to the iterative centres. The effects of different group centre choices on the group-galaxy cross-correlation concern mainly the 1-halo regime at $r \leq 1 h^{-1}$ Mpc, and the correlation functions converge on larger scales (Yang et al. 2005).

The halo mass of GAMA groups was found to be tightly correlated with the group total luminosity (Han et al. 2015; Viola et al. 2015; Rana et al. 2022), based on using stacked weak lensing measurements to determine the mass distribution of the GAMA groups. The halo mass of groups is related to the r -band luminosity L_{grp} via

$$M_h = M_p \left(\frac{L_{\text{grp}}}{L_0} \right)^\alpha, \quad (6)$$

where $L_0 = 2 \times 10^{11} h^{-2} L_\odot$, $\log_{10}(M_p/h^{-1} M_\odot) = 13.48 - 0.08 \pm 0.12$, and $\alpha = 1.08 + 0.01 \pm 0.22$ (Han et al. 2015). The group halo mass can be used to compute the expected mean group bias, as shown in Section 5.3, where we also consider alternative mass calibrations. It should be noted that in these works, only groups with three members or above are used. Robotham et al. (2011) showed that the mass function is noisier, but not biased when including two-member groups. Although these systems are individually of low reliability, this aspect should be allowed for by the mock catalogue, allowing us to gain the statistical advantage of using a larger sample. The luminosity is computed from the apparent r -band magnitude:

$$-2.5 \log(L/L_\odot) = m - K(z) - 5 \log(d_L) - 25 - M_\odot, \quad (7)$$

where $K(z)$ the k -correction up to $z = 0$ (KCORR_Z00), d_L is the luminosity distance, and $M_\odot = 4.67$ is the r -band absolute magnitude of the sun. The luminosity distance is expressed with unit h^{-1} Mpc so that the luminosity has units of $h^{-2} L_\odot$.

The total luminosity of the group is computed in Robotham et al. (2011) via

$$L_{\text{FoF}} = B L_{\text{ob}} \frac{\int_{-30}^{-14} 10^{-0.4M_r} \phi_{\text{GAMA}}(M_r) dM_r}{\int_{-30}^{M_{r-\text{lim}}} 10^{-0.4M_r} \phi_{\text{GAMA}}(M_r) dM_r}, \quad (8)$$

where L_{ob} is the total observed luminosity in the r_{AB} band, $B = 1.04$ is the correction for median unbiased estimate for $N \geq 5$ groups, and $M_{r-\text{lim}}$ is the absolute magnitude limit of the group depending on the redshift z . ϕ_{GAMA} is the luminosity function defined in Robotham et al. (2011). The luminosity function at the faint end for GAMA galaxies is well approximated by $\phi \propto L^{-1} \exp(-L/L_*)$ (Loveday et al. 2012). Thus in practice we take the simpler approach of estimating a total group luminosity by scaling the observed luminosity by a redshift-dependent correction factor $\exp(z^2/z_*^2)$ with $z_* = 0.33$, where z is the mean redshift of the group members. This correction factor has been checked using the G3C groups to produce a total luminosity consistent with the official TotFluxProxy.

The total stellar mass is another proxy for the total group mass. We take the StellarMassesv19 DMU from Taylor et al. (2011), where stellar population synthesis is used to model the optical photometry of the GAMA galaxies. Because the modelling uses rest frame luminosities, which depends on distance, the stellar mass is expressed in units of $h^{-2} M_\odot$.² Furthermore, for each group, we correct the total stellar mass by the same redshift dependent factor as the total luminosity. Notice that we do not apply the fluxscale correction here, which accounts for the missing flux from matched aperture photometry, because our results do not rely on the absolute stellar mass of the groups. This correction therefore does not affect our primary aim of splitting the groups into a few bins based on their ranking in mass.

The calibration of the total stellar mass and the halo mass from weak lensing of the GAMA groups is shown in Fig. 2 for the official G3C groups from the G3CFoFGroupv09 DMU (dashed line) and the group catalogue used in this work (solid line). The contours show 95, 50, and 20 per cent of the total sample, and are highly consistent between the two group catalogues. We choose to divide groups into three stellar mass bins based on percentiles: the low mass (LM) bin consists of the least massive 40 per cent groups, the medium mass (MM) bin corresponds to the middle 50 per cent, and the high mass (HM) bin contains the most massive 10 per cent. The signal-to-noise of high mass haloes is expected to be high, despite the low number in the HM bin.

2.3 Mocks

We include mock catalogues for two reasons: (1) to validate the RSD models and assess the bias on the recovered growth rate, and (2) to quantify the impact of cosmic variance via the construction of covariance matrices. We used 25 realizations of a light-cone mock catalogue based on the GALFORM semi-analytical galaxy formation (Gonzalez-Perez et al. 2014). The catalogue exploits the Millennium Simulation (Boylan-Kolchin et al. 2009) with the WMAP7 cosmology. These mocks can be obtained from the Durham hosted Virgo-Millennium Database1³ (Lemson & Virgo Consortium

²Notice that this is only approximately true, because the stellar mass to light ratio, M/L , which is used to obtain the stellar mass, depends on age and is therefore specific to the choice of h . The stellar mass used here assumes $h = 0.72$.

³http://virgodb.dur.ac.uk:8080/MyMillennium/Help?page=databases/gama_v1/lc_multi_gonzalez2014a

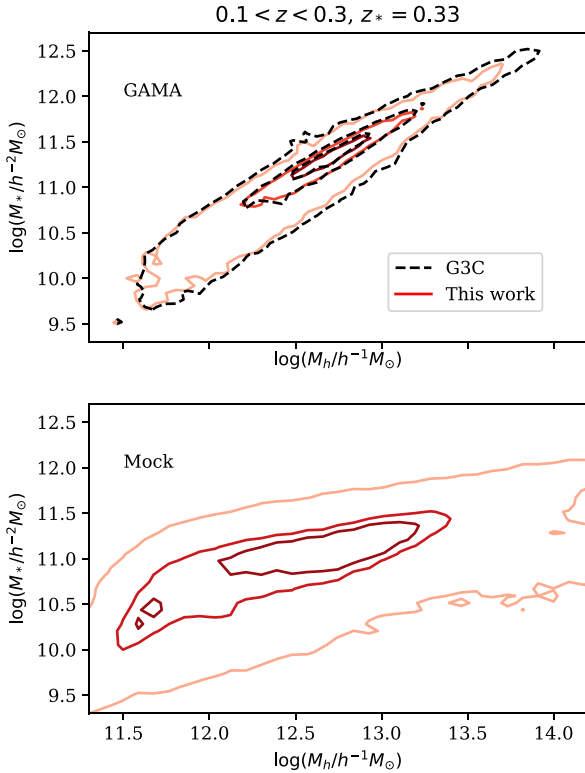


Figure 2. Upper panel: Correlation between the stellar mass, corrected to total by a factor $\exp(z^2/z_*^2)$ and the halo mass estimate derived from the total group luminosity mass proxy, together with a lensing-based absolute calibration from Han et al. (2015), for the GAMA groups with two or more members between redshifts $0.1 < z < 0.3$. The contours denote 95, 50, and 20 per cent of the total sample. The solid lines show the groups used in this work using the group finder algorithm in Treyer et al. (2018), and the dashed lines show the official G3C groups (Robotham et al. 2011). Lower panel: The same relation for the mock catalogue. In this case, M_h is not estimated from the luminosity, but directly taken as the arithmetic mean host halo mass of the group member. The difference in the distributions indicates that the stellar populations in the mock data are not entirely realistic, but it also warns us that the exact values of halo mass corresponding to the different GAMA group subsets must be treated empirically, and should not be treated as being known precisely.

2006). For more details regarding the mock catalogue, see Farrow et al. (2015). By equation (2), the fiducial value of growth rate at the mean redshift of the mocks, $z = 0.195$, is $f_{\text{fid}} = 0.593$. The light-cone is constructed using the methods in Merson et al. (2013), where, given an observer, the galaxy is placed at the epoch where it first enters the past light-cone of the observer. The galaxy trajectories are interpolated between snapshots. Each mock covers the five GAMA fields with the SDSS r -band apparent magnitude $\text{SDSS}_r\text{-obs_app} < 21$, and $z < 0.9$.

We use galaxies in the G09, G12, and G15 fields and apply the same selection in redshifts $0.1 < z < 0.3$ and the apparent r -band magnitude cut $\text{SDSS}_r\text{-obs_app} < 19.8$. We also apply the same survey mask generated using the random catalogue. The masked areas are obtained by binning random galaxies in each field with an average of ~ 2000 counts in each bin. Pixels with counts smaller than five times the Poisson noise are masked. The total masked area in the three fields is about 0.14 deg^2 . Because the mock redshift distribution is not matched exactly with GAMA data and random (see Fig. 3), we create a random catalogue for these mocks by down-sampling the

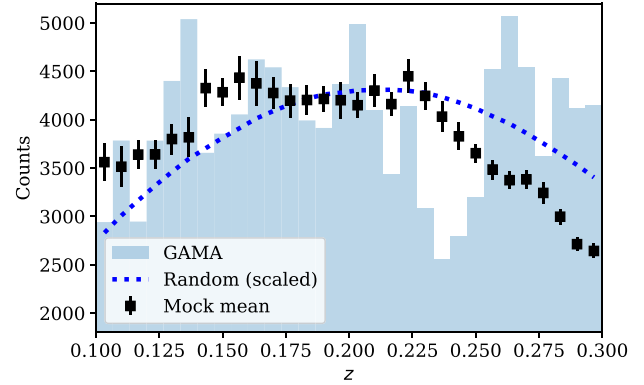


Figure 3. The mean redshift distribution of the 25 GAMA mocks (square) is offset from that of the random sample (dotted line). A random catalogue is created for the mocks to have matched redshift distribution as the mock mean. The redshift distribution of the GAMA galaxy sample is also shown (histogram) for comparison.

random catalogue for the GAMA data, such that the $n(z)$ matches the mean of 25 mocks.

The red and blue subsamples for the mean of the mocks are separated by the empirical line given by

$$g - i = 0.46 + 3.2z, \quad (9)$$

as shown in the lower panel of Fig. 1. The line is chosen to go through the green valley of the mock galaxy $g - i$ colour. The GAMA galaxies have a more concentrated red sequence overlapping with an extended blue population, without a distinct green valley in between. On the contrary, the mocks have a broader red population which is well separated from the blue population by a green valley. Since the mock catalogues have more distinctive separation for the two populations, we find the corresponding colour cut in the GAMA data by matching red and blue fractions in the two catalogues for 20 redshift bins in $0.1 < z < 0.3$. The cut is smoothed by fitting a second-order polynomial, as shown in the upper panel of Fig. 1.

The contamination of the red and blue sub-samples in the GAMA data resulting from the colour cut is quantified in the following way: for each redshift bin, the red and blue sub-samples are fitted by a double Gaussian. It is a reasonable fit except for the green valley in the mocks, as shown in Fig. 4. Given a colour cut, the contamination of the red sub-sample is defined as the area under the blue Gaussian over the area under the red Gaussian, and similarly for the contamination of the blue sub-sample. Clearly, GAMA data contain a contaminated red sample and a pure blue sample. Therefore, we create a contaminated red sub-sample using the mock catalogues by placing the mock colour cut such that extra blue galaxies are included with the same level of contamination as GAMA data. The contaminated red cut in the mocks (see Fig. 1) is smoothed by fitting a quadratic polynomial of the form

$$g - i = 2.43z^2 + 1.55z + 0.388. \quad (10)$$

For mock groups, the stellar mass is computed by the sum of `diskstellarmass` and `bulgemass` of all group members, and corrected by the same redshift-dependent factor as the data. We do not estimate the group halo mass from the same mass–luminosity relation in equation (6). Instead, we use the host halo mass of the mock galaxy directly. Because some haloes contain more than one galaxy, for each group, we test the largest, the arithmetic mean, and the median halo mass of the group member, and find that they give

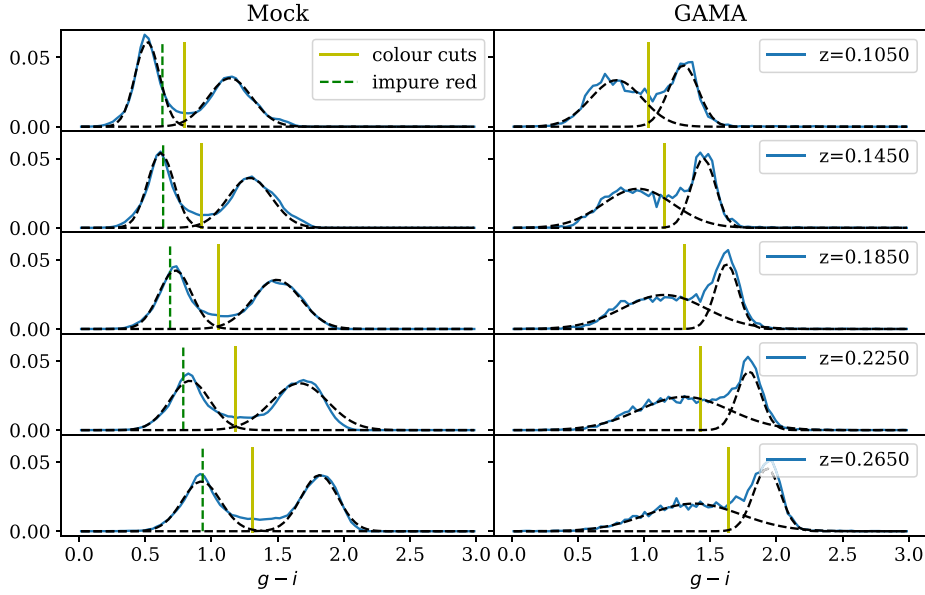


Figure 4. The $g - i$ colour distribution of GAMA and 25 mocks combined at 5 of the 20 redshift bins in $0.1 < z < 0.3$. The black dashed lines are double Gaussian fits to the distributions, characterizing the blue and red populations. The yellow cuts show a linear cut in the green valley in the mocks, and the corresponding cuts in GAMA which give the same red and blue fraction. The green dotted cuts in mock is the cut for the impure red sample.

similar results. We also test using the sum of unique host haloes in the group. This increases the total group halo mass in the lower mass end, but does not affect the higher mass end. The stellar–halo mass relation of the groups using the total stellar mass and the arithmetic mean host halo mass of the group members is shown in the lower panel of Fig. 2. It is clear that the mocks show a much larger scatter in the $M_h - M_*$ plane and the slope is smaller compared to data, i.e. at fixed stellar mass, the halo mass is larger. The total stellar mass of the mock groups is also smaller by about 0.5 dex compared to data. The clear difference between data and the mocks shows that estimating the halo mass from luminosity using equation (6) is not very reliable. The luminosity is itself strongly correlated with stellar mass via the luminosity–mass relation, thus the upper panel of Fig. 2 does not show the true scatter of M_h at fixed M_* faithfully (or vice versa). A comparison between the group and the halo catalogue in the GAMA mocks reveals that the 2-member groups have low fidelity, also discussed in Robotham et al. (2011). This again emphasizes the importance of using a consistent group finder algorithm between the GAMA data and the mock catalogues. The mock group catalogues are separated into three stellar mass bins based on the 40, 50, and 10 per cent percentiles as measured in the data.

3 CROSS-CORRELATION MEASUREMENTS

3.1 Correlation statistics

We estimate the 2-point correlation function by counting pairs of galaxies and randoms using the Davis–Peebles estimator (Davis & Peebles 1983):

$$\hat{\xi}(r_p, \pi) = \frac{D_1 D_2}{D_1 R_2} - 1, \quad (11)$$

where the subscript $i = 1, 2$ denotes the two samples to be correlated (in case of autocorrelation, the same sample), D_i denotes data, and R_i denotes the corresponding random points. Each term in the equation (e.g. $D_1 D_2$) is the normalized pair count between data and

(or) random points, measured in bins of the pair separation (r_p, π) defined below. For two objects located at $\mathbf{s}_1$ and $\mathbf{s}_2$, their separation is given by $\mathbf{s} = \mathbf{s}_1 - \mathbf{s}_2$. The line of sight is defined along the mean position of the pair, $\mathbf{r} = (\mathbf{s}_1 + \mathbf{s}_2)/2$. One can then decompose the separation into components parallel and perpendicular to the line of sight:

$$\pi = \frac{\mathbf{s} \cdot \mathbf{r}}{|\mathbf{r}|}; \quad r_p = \sqrt{s^2 - \pi^2}. \quad (12)$$

The random catalogue R captures various properties of the actual data, such as the survey mask and the sample redshift distribution $n(z)$, but has no spatial correlation. Thus, these estimators essentially measure the excess clustering of the data points compared to a random distribution. For the red and blue galaxy samples, the random $n(z)$ is modulated by the redshift-dependent red and blue fraction respectively (see Section 2.1).

For the case of autocorrelations, the Landy–Szalay estimator (Landy & Szalay 1993) is known to be superior to the Davis–Peebles approach, and there is a natural generalization to cross-correlation:

$$\hat{\xi}(r_p, \pi) = \frac{D_1 D_2 - D_1 R_2 - D_2 R_1 + R_1 R_2}{R_1 R_2}. \quad (13)$$

However, implementing this estimator would require a random catalogue for galaxy groups in different mass ranges, and we prefer to avoid this complication. In contrast, the Davis–Peebles estimator requires a random catalogue for only one of the populations being correlated. Both Mohammad et al. (2016) and Riggs et al. (2021) estimated cross-correlations using a form of Landy–Szalay where R_2 was replaced by R_1 , but this has no justification, and will yield incorrect results when the selection functions of the two tracers are very different. We measured the galaxy autocorrelations using both estimators, and found negligible difference for our sample. Throughout the analysis, the size of random catalogue used is 20 times that of data.

The 2D correlation functions are measured out to a maximum scale of $40 h^{-1} \text{ Mpc}$ for both the r_p and π directions in bins of

$1 h^{-1}$ Mpc. This maximum scale is chosen due to the limited volume of the GAMA survey. It may appear to be a concern that at this scale perturbation theory starts to break down (thus the scale is often chosen as the cut-off scales for larger data sets). However, we shall show later that empirical models based on linear theory can still give relatively unbiased results below this scale. For the halo streaming model, which we will elaborate in Section 4.2, we use a 1-halo template to absorb the deviation of non-linear clustering from the perturbative 2-halo term. Although in principle, r_p and π should be measured in the range $[-40, 40] h^{-1}$ Mpc, in practice, pair counts in the positive and negative bins are combined and our correlation functions have two mirror planes of symmetry. This is the standard practice for r_p , because the correlation function is symmetric around the transverse direction. Along the line of sight, positive and negative π measurements can in fact be distinctive in cross-correlations due to secondary gravitational effects. For example, gravitational redshift can give rise to a non-vanishing dipole between two samples that differ significantly in mass (Wojtak, Hansen & Hjorth 2011; Bonvin, Hui & Gaztañaga 2014; Cai et al. 2017; Beutler & Dio 2020). However, given the size of the sample, we shall not investigate this issue in this study. The combination of the π bins also improves the signal-to-noise ratio for our measurements.

3.2 Results

In the following analysis, we will refer to the group subsamples as LM, MM, and HM. We thus have six configurations for the group-galaxy cross-correlations: LMred, MMred, HMred, LMblue, MMblue, and HMblue. In addition, we also measure the red and blue galaxy autocorrelations. The inclusion of galaxy autocorrelations in the analysis helps in breaking the near degeneracy between b_{gal} and b_{grp} . Ideally, one would also include the autocorrelations for the group catalogue. These are excluded here: as mentioned above, we do not construct random versions of the group catalogue.

Fig. 5 shows the red and blue galaxy autocorrelations (top row) and the cross-correlation functions for the six configurations (lower three rows) measured from the GAMA data (first and third column), and the corresponding mocks average (second and forth column). The first noticeable feature is that the red configurations (left two columns) have larger clustering signals compared to the blue ones (right two columns). This is most obvious in the two galaxy autocorrelations. On relatively large scales, the squashing is also stronger in the blue configurations. This effect is controlled by the Kaiser distortion parameter $\beta = fb$ (see Section 4.1 below), and is expected to be stronger for samples with a smaller bias b (vice versa), given that the growth rate f is fixed. The fact that red galaxies have a larger galaxy bias compared to blue galaxies implies that red galaxies are preferentially associated with more massive dark matter haloes, in agreement with other studies (e.g. Guo et al. 2014; Mandelbaum et al. 2016; Bilicki et al. 2021). On smaller scales, the red configurations also show a much more prominent FoG signal compared to the blue ones. This is intuitively sensible because more massive haloes are associated with a larger velocity dispersion. Comparing the correlation functions across different groups for a specific galaxy selection, we see a similar trend on both small and large scales: with the increase in group mass, a larger clustering amplitude, group bias, and FoG effect are observed. Notice the high signal-to-noise ratio at small scales in the HM groups, despite that the sample size in this mass range is only about 1/4 of the other mass ranges. These observations confirm that the identification of the galaxy groups, as well as the separation of group masses based on the effective

halo mass (total luminosity) are successful for the purpose of this study.

The agreement between the mock average and the GAMA data is good in general – the same trends in galaxy colour and group mass are captured. In regions where r_p is close to zero, the mock average seem to produce weaker clustering compared to the actual data in the blue configurations. The match in the red configuration, on the other hand, is excellent. The mock contaminated red sample is also shown as dotted contours (second column). The inclusion of extra blue galaxies has the effect of slightly reducing the overall amplitude in this contaminated sample compared to the pure red sample. On larger scales, the signal in data is noise and cosmic variance dominated. This is most noticeable in the LM subsamples, where the signal greatly exceeds the mock average on $r \geq 20 h^{-1}$ Mpc. Inspecting the measurement in each mock sample, this level of fluctuation in the data is expected. It should be noted that cosmology adopted in the mock catalogue is $\Omega_m = 0.27$, which is lower than the current constraint from *Planck*, $\Omega_m = 0.315 \pm 0.007$ (Planck Collaboration VI 2020). Thus, one may expect some difference in clustering between the mock average and the data. However, given the noise in the GAMA data, a percent-level shift in the growth rate $f \propto \Omega_m^{0.55}$ is hard to discern. It is also found in Farrow et al. (2015) (e.g. their fig. 8) that the mock can capture similar clustering trends as the GAMA data, when split into bins of redshift and stellar mass. Notice that there are significant deviations at small scales ($r_p < 1 h^{-1}$ Mpc) in the shape of the projected correlation functions, but these scales are not explored in this analysis.

4 RSD MODELS

4.1 Quasi-linear dispersion model

To describe the RSD in galaxy density field, one can extend equation (3) by including the galaxy bias b :

$$P_g^s(k, \mu) = b^2 P^r(k) (1 + \beta \mu^2)^2, \quad (14)$$

where $\beta = fb$ is often referred to as the distortion parameter. The above formalism is valid for galaxy autocorrelation, but it is straightforward to generalize to cross-correlation:

$$P_c^s(k, \mu) = \frac{b_{\text{gal}}^2}{b_{12}} (1 + \beta_{\text{gal}} \mu^2) (1 + b_{12} \beta_{\text{gal}} \mu^2) P^r(k), \quad (15)$$

where b_{gal} is the galaxy bias, and b_{12} is the ratio between galaxy and group bias:

$$b_{12} \equiv b_{\text{gal}}/b_{\text{grp}}. \quad (16)$$

The 2-point correlation function is the Fourier transform of the power spectrum. It is convenient to express the correlation functions in terms of Legendre polynomials $\mathcal{P}_\ell(\mu)$ with $\ell = 0, 2, 4$ (Hamilton 1992):

$$\xi_g^s = \xi_0(r) \mathcal{P}_0(\mu) + \xi_2(r) \mathcal{P}_2(\mu) + \xi_4(r) \mathcal{P}_4(\mu). \quad (17)$$

In this expression, the coefficients of the Legendre polynomials, $\xi_0(r)$, $\xi_2(r)$, and $\xi_4(r)$ are referred to as the monopole, quadrupole, and hexadecapole. In linear theory, only even modes are present up to the forth order because of the RSD effect modifies the power spectrum by the factor $(1 + \beta \mu^2)^2$. The specific form of these multipoles are computed by Hamilton (1992) for autocorrelation, and Mohammad et al. (2016) for cross-correlation. We summarize these formulae in Appendix A.

The FoG effect is accounted for by a convolution of the correlation function with some distribution of the non-linear random peculiar velocity along the line of sight (Peacock & Dodds 1994). N -body

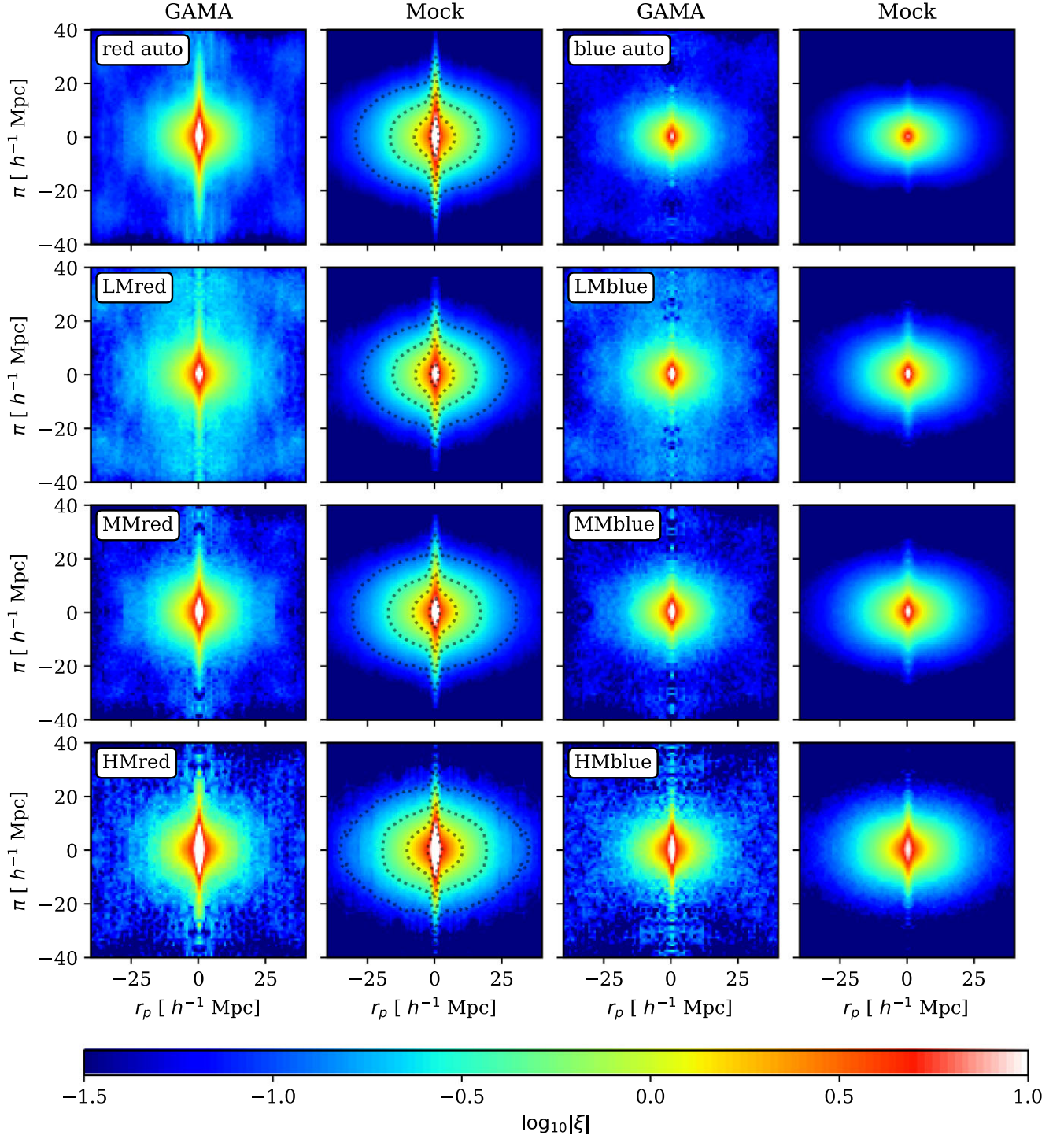


Figure 5. False colour images of autocorrelation and cross-correlation functions in redshift space for the actual GAMA data and the corresponding average over the set of 25 GAMA mocks. r_p denotes transverse separation; π is radial separation. LM, MM, and HM denote the three group mass bins. A number of trends are apparent: both the bias (amplitude of clustering) and the small-scale Finger of God (FoG) dispersion increase with group mass, and are larger for red galaxies than for blue. The mock contaminated red sample is shown as dotted contours on the second column, with $\log_{10}|\xi| = \{-1.0, -0.5, 0, 0.7\}$.

simulations show that the actual distribution is non-Gaussian (e.g. Sheth 1996; Scoccimarro 2004; Cuesta-Lazaro et al. 2020). Thus, we adopt:

$$D(\pi) = \frac{1}{\sqrt{2}\sigma_{12}} \exp\left(-\sqrt{2}H_0\pi/\sigma_{12}\right), \quad (18)$$

where σ_{12} is the pairwise velocity dispersion. In Fourier space, this function takes the form of a Lorentzian function, $\tilde{D}(k\mu) = [1 + (k\mu\sigma_{12})^2/2]^{-1}$, which damps the high- k modes of the anisotropic power spectrum.

At quasi-linear scales, non-linearity may introduce systematic biases in the inferred cosmological parameters (de la Torre & Guzzo 2012). There are multiple challenges in extending the model beyond

linear regime. There, the peculiar velocities can be large, and the formalism described breaks down at linear order. Non-linearities alter the small-scale shape of the matter power spectrum and correlate the density and velocity fluctuations. Accounting for these effects requires higher order expansion in the Perturbation Theory and the inclusion of the velocity spectrum, $P_{\theta\theta}(k)$, and the density–velocity cross spectrum, $P_{\delta\theta}(k)$, e.g. the TNS model by Taruya, Nishimichi & Saito (2010). Galaxy bias can also be non-linear and stochastic on small scales (Dekel & Lahav 1999). Furthermore, the approximate velocity dispersion in equation (18) fails to fit autocorrelation data on the smallest scales. More elaborate velocity distributions are proposed by e.g. Reid & White (2011); Zu & Weinberg (2013); Bianchi, Chiesa & Guzzo (2015) based on simulations.

One simple approach is the replacement of the linear power spectrum in the linear Kaiser model (equation 14) by the non-linear power spectrum. This is reasonable because the redshift space power spectrum should match that in real space at $\mu = 0$. Blake et al. (2011) showed that this combination is actually among the best-performing RSD models when fitting down to $k_{\max} = 0.2 \, h \, \text{Mpc}^{-1}$ with fixed cosmology.⁴ For this model, we adopt the non-linear power spectrum from HALOFIT (Smith et al. 2003; Takahashi et al. 2012). In the non-linear regime, we should in principle allow for a scale-dependent bias. But in practice it is a good approximation to assume that the non-linear galaxy and matter power spectra are in a constant ratio (see Camacho et al. 2019). In the following analysis, we refer to this model as the ‘quasilinear dispersion’ (QD) model.

4.2 RSD in the halo model

The main deficiency of the QD model is that it does not address post-linear couplings between density and velocity, which will modify the simple Kaiser angular anisotropy. There is an extensive literature of attempts to improve such modelling, based on various forms of perturbation theory. The model of Taruya et al. (2010) is widely used, although more recent efforts have concentrated on the Effective Field Theory approach. This adds additional terms dictated by symmetry in a way that can also capture bias effects, including non-linearity and non-locality (e.g. Carrasco, Hertzberg & Senatore 2012; Senatore 2015; d’Amico et al. 2020). These results are impressive, but have the limitation that they are presented in Fourier space and are not reliable beyond $k \simeq 0.3 \, h \, \text{Mpc}^{-1}$. For a robust prediction of correlation functions, we need a formalism that still behaves correctly in the large- k limit.

For this reason, we have developed a model that seeks to access the highly non-linear regime by using the halo model. In real space, this involves correlations that count pairs of galaxies in the same halo or in different haloes:

$$\xi(r) = \xi_{1h}(r) + \xi_{2h}(r). \quad (19)$$

The 1-halo term is determined by the form of the halo density profile, and the 2-halo term is close to a linearly biased version of the matter two-point function. The bias in turn is determined by the halo occupation number, $N(M)$, of galaxies in haloes as a function of their mass. This halo model has proved a highly effective way to understand the relation between the clustering of galaxies and of mass (Peacock & Smith 2000; Seljak 2000; Cooray & Sheth 2002), and for the case of dark matter alone has led to the highly precise HALOFIT framework (Smith et al. 2003; Takahashi et al. 2012).

⁴ Although it should be noted that if the model could introduce bias to Ω_m if the cosmology is not fixed, as shown in Parkinson et al. (2012).

The halo-model separation into two independent pair contributions must also apply for the redshift-space correlations, namely

$$\xi(r_p, \pi) = \xi_{1h}(r_p, \pi) + \xi_{2h}(r_p, \pi), \quad (20)$$

but it should be clear from the outset that the 1-halo and 2-halo contributions would be expected to have rather different anisotropy signals. The characteristic quadrupole plus hexadecapole Kaiser distortion arises from the coherent component of the velocity field, and this will apply to the 2-halo term only, since pairs from within the same halo are unaffected by bulk motion of the halo. This redshift-space decomposition using the halo model was advocated by Hand et al. (2017), who invested much effort in trying to predict the two distinct components using perturbation theory. Our work bears some resemblance to their approach, with two distinct differences: we work directly in configuration space, and we base the 1-halo term on empirical simulation results, rather than attempting to calculate it a priori.

A particular point to clarify in this decomposition is the treatment of Fingers of God. Random motions within a halo are treated in the dispersion model by a radial convolution – but in fact the appropriate convolution will be different for the 1-halo and 2-halo terms. The main reason for this is that the 1-halo and 2-halo terms weight contributions as a function of halo mass differently, with a higher weight given to high-mass haloes in the 1-halo term (see e.g. equations 8 and 10 of Seljak 2000). Since the pairwise dispersion σ_{12} increases with halo mass, we expect larger FoG effects to apply to the 1-halo term. This is further complicated by the existence of central and satellite galaxies, since the weighting of these is different in the 1-halo and 2-halo terms. For example, suppose each halo contains either just a single central or one central and one satellite, where the velocity dispersion of satellites is σ . The 1-halo contribution must pair a central with a satellite, so the pairwise dispersion is σ . But the 2-halo term can also pair centrals with centrals (assumed to have negligible pairwise dispersion – although Reid et al. 2014 showed that the actual pairwise velocity could be up to 30 per cent of σ) and satellites with satellites (pairwise dispersion $\sqrt{2}\sigma$), so the average rms pairwise dispersion depends on the fraction of haloes that contain a satellite. If most haloes are central-only (as in BOSS CMASS, for example), the pairwise dispersion for the 2-halo term will be $\ll \sigma$. In the opposite direction, one can argue that the velocity field of haloes will contain some stochastic component in addition to the coherent velocities that generate the Kaiser distortion.

With this perspective, an improved simple model for the cross-power between tracers a and b would be as follows:

$$P_{ab}(k, \mu) = P_{1h}(k) D_1(k\mu) + b_a b_b P_{\text{lin}}(k) (1 + \beta_a \mu^2)(1 + \beta_b \mu^2) D_2(k\mu). \quad (21)$$

Leaving aside the 1-halo term for the moment, one way in which we can seek to improve this expression further is in terms of quasi-linear effects on the 2-halo term. A first requirement is that the real-space spectrum (at $\mu = 0$) should have the full non-linear form. When discussing the dispersion model, we achieved this by replacing P_{lin} by the non-linear spectrum. In the halo model, we should not do this, since the 2-halo term in real space is close to linear theory, and the 1-halo term supplies most of the non-linear corrections (Smith et al. 2003). We do however adopt the HALOFIT 2-halo term, with scale-independent bias, as the best model for the real-space 2-halo term.

The next step is to seek improvement in the density–velocity coupling that leads to the Kaiser distortion factors. An attractive approach here is the streaming model (e.g. Fisher 1995; Vlah,

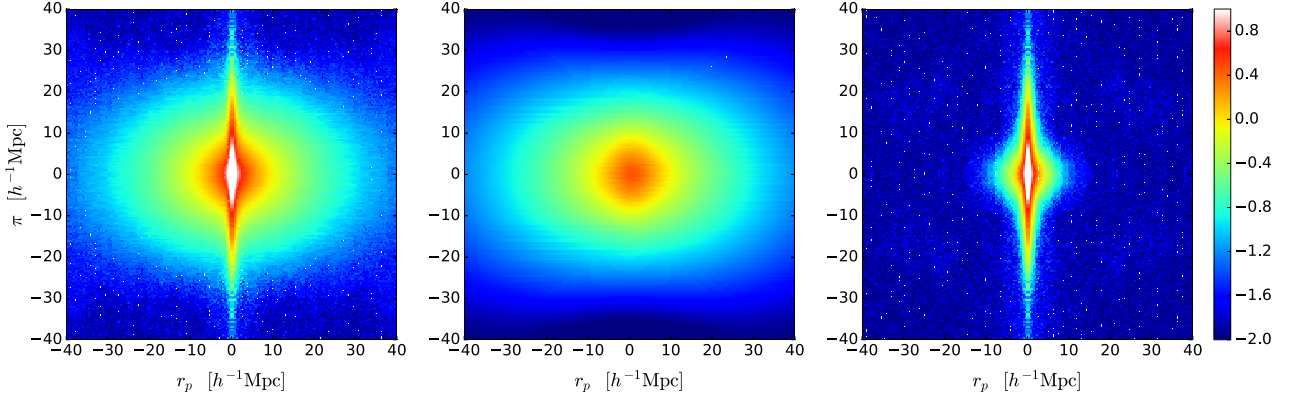


Figure 6. Illustrating the decomposition of the measured mock correlation data (left-hand panel) into a 2-halo fit (middle panel) and an empirical 1-halo term in the form of the residual of the fit (right-hand panel), for the particular case of red-MM cross-correlation. The 2-halo term is computed using the streaming model, and is matched to the data at radii $r > 10 h^{-1}$ Mpc, with the additional criterion that $r_p > 3 h^{-1}$ Mpc.

Castorina & White 2016), in which we consider the quasilinear relative velocity distribution as a function of pair separation, and use this to transform to redshift space while exactly conserving pair counts. The details of the construction of this model are given in Appendix B. As with the linear model for the 2-halo term, there are three main free parameters, the tracer biases and the growth rate: (b_a, b_b, f) . This assumes that the mass power spectrum is known exactly, whereas it depends on all fundamental Λ CDM parameters. The main variation of the power is $\propto \sigma_8^2$, so it is common to factor out this degree of freedom and take the main RSD parameters to be $(b_a \sigma_8, b_b \sigma_8, f \sigma_8)$. However, there is a weaker further dependence on σ_8 when we adopt the HALOFIT prediction of the 2-halo matter power spectrum, rather than taking this to be pure linear theory (although the difference is not important in practice).

In addition to these three main RSD parameters, we have parameters connected to the FoG damping. As described earlier, it is conventional to model FoG effects by radial convolution, taking the velocity PDF to be a Lorentzian and using a velocity dispersion as the single free parameter. But in the present context, it is important to be clear that the empirical evidence for the Lorentzian form comes mainly from the 1-halo term. This is $D_1(k_z)$ in equation (21); but we want the effect on the 2-halo term, $D_2(k_z)$. We have argued that this will be characterized by a different dispersion, but in addition there is no strong reason to assume it will have a Lorentzian form. As a more general alternative, we considered a modified Lorentzian:

$$D_2(k\mu) = (1 + (k\mu\sigma_{12})^2/2\gamma)^{-\gamma} \quad (22)$$

and experimented with different values of γ . But in practice the results were rather insensitive to the choice of this parameter, so we retained the Lorentzian $\gamma = 1$. It is shown in e.g. Scoccimarro (2004), Bianchi et al. (2015), and Cuesta-Lazaro et al. (2020) that the shape of the PDF is only relevant where the correlation function changes significantly over a scale comparable to the width of the smoothing function. But the issue of the exact form of the PDF for FoG corrections to the 2-halo term is a problem that merits further study.

In summary, we therefore have two models with similar real-space correlations, but different degrees of RSD: (1) QD: quasi-linear dispersion model and (2) HS: halo streaming model. Both of these converge to the linear Kaiser model on large scales, so what is of interest is the smallest scale to which their predictions are reliable. We will assess these by comparison with mock data.

4.2.1 The 1-halo term

The real-space 1-halo term can in principle be computed in the usual halo model framework, given the occupation numbers for the tracers and the halo radial profile. But there is also a case for taking an empirical approach, given that the real-space correlations are in principle observable directly, in a manner free of RSD effects, via the projected correlation function $w_p(r_p)$. One might for example model the real-space 1-halo term by a power-law of free amplitude and slope, or via an NFW profile.

But whatever approach is taken in real space, there is then the question of how the 1-halo term appears in redshift space. As described above, the simplest approach is to assume that the transition to redshift space consists of a radial convolution with a single FoG function. However, it is not hard to see that this must be an oversimplification. The 1-halo term arises from random orbital velocities within the halo, but the velocity dispersion is unlikely to be constant. If for example we consider the case of isotropic orbits, then the dispersion would need to fall to zero at the virial radius of the halo, beyond which the density is assumed to vanish.

Here, we address this concern directly by using the mocks. Given a hypothesis for the 2-halo term, we can subtract the 2-halo prediction from the mock data to obtain an empirical $\xi_{1h}(r_p, \pi)$ that sums with the 2-halo term to give exactly the mock data (specifically, we apply this approach to the average of all the mocks). The 2-halo term can be deduced by fitting to the mock data in a regime where we assume the 1-halo contribution to be negligible. The exact cuts adopted in the process are not critical; in practice, we chose to match to the data at radii $r > 10 h^{-1}$ Mpc, with the additional criterion that $r_p > 3 h^{-1}$ Mpc. The operation of this procedure is illustrated in Fig. 6. The resulting residual 1-halo term is clearly well localized near the origin, and indeed it can be seen that the RSD effects in the 1-halo term are complicated, with the FoG effect being largest at $r_p = 0$, whereas the function appears more isotropic close to its outer limit at $r_p \simeq 5 h^{-1}$ Mpc. This interesting behaviour is clearly worthy of being modelled in detail, but we shall not do that here.

We now have a decomposition of the redshift-space correlations that by construction exactly matches the average of the mocks. However, each mock realization will be different, as will be the real data, so can these different data sets be fitted in this framework? The 2-halo term is already parametrized, and these parameters can be varied for any given data set. But the 1-halo term must also have some variation. Our approach is to assume that the mocks are sufficiently

realistic that the effective 1-halo term in any given case will be close to the mock average, and that the difference can be captured by two nuisance parameters:

$$\xi_{1h}(r_p, \pi) \rightarrow \alpha \xi_{1h}(r_p, \eta \pi). \quad (23)$$

In other words, assume that we have roughly the right functional form, but that the amplitude may be off (scale by α), and that the FoG strength may be off (stretch in the radial direction by $1/\eta$). Physically, the amplitude parameter α can be related to the way in which galaxies populate haloes: there may be different numbers of satellite galaxies in a given halo compared to the mock, leading to a different small-scale clustering signal. We have previously seen that an empirical rescaling of the 1-halo amplitude can yield an accurate fit to correlation data in real space (Hang et al. 2021). The η parameter attempts to capture the velocity dispersion of the galaxies in the halo: a smaller η produces a larger FoG effect. Again, this can be understood in terms of an uncertain halo occupation, which can alter the mean mass of the haloes that contribute the 1-halo term.

As we show below, this approach is able to succeed in matching the individual mock realizations, and so we see no reason not to apply the same model to the real data. We emphasize that we do not need to assume that the mocks are completely realistic, as long as they are qualitatively similar to reality. The reliability of this approach can be judged by whether or not the fitted values of α and η are close to unity (as indeed turns out to be the case). We only consider this minimal set of two empirical nuisance parameters in the current analysis; for forthcoming large data sets of higher statistical precision, more parameters may be required in order to make the model acceptably accurate. Eventually, we will need to validate the model by deriving 1-halo templates from a given set of mocks and showing that they can fit data derived from mocks produced according to different assumptions. We intend to pursue this high-precision robustness test in a future study.

4.3 Fitting methodology

4.3.1 Covariance matrix and likelihood inference

In the above discussion, we have not been explicit about exactly what it means to fit the averaged mock data. In principle, one would like to have an understanding of the errors on the data, so that the likelihood can be computed as a figure of merit that is used to optimize the fit. For an individual data set, this can be done in the standard way by using an ensemble of mocks to estimate the covariance matrix of the data, and then appealing to the central limit theorem to compute the likelihood in the Gaussian approximation. For fitting the stacked mocks, the appropriate covariance matrix is less obvious, but in any case it is less important to have a likelihood in that case, where the aim is simply to estimate a 1-halo contribution as a basis for further modelling. We are not interested in placing errors on the best-fitting parameters of the 2-halo term, for which a likelihood would be required. In practice, therefore, we took the simple approach of seeking a least-squares fit in $\ln(1 + \xi)$ to the mock average. The exact figure of merit chosen is unimportant as regards the 1-halo residual.

The covariance matrix for a single data realization is most often estimated in one of two ways: either directly via the scatter over a number of mock realizations, or via Jackknife resampling of a single realization. Both of these approaches have their limitations, but the best strategy is when they are combined: an expanded set of mock realizations is created by Jackknife resampling of each one, yielding an improved estimate of the covariance matrix (Alam et al. 2021b).

For a data vector with component x_i and a model vector y_i , where $i = 1, \dots, m$, the χ^2 is defined as

$$\chi^2 = \sum_{i,j} [x_i - y_i] C_{ij}^{-1} [x_j - y_j]. \quad (24)$$

In the above equation, C_{ij} is the covariance matrix, estimated from N independent realizations of mock data:

$$\hat{C}_{ij} = \frac{1}{N-1} \sum_{k=1}^N [x_i^k - \langle x_i^k \rangle] [x_j^k - \langle x_j^k \rangle]. \quad (25)$$

Given a model with p parameters, there are $m - p$ degrees of freedom in the χ^2 -fitting. Due to the small number of mocks, we apply Jackknife re-sampling on the mocks by dividing each survey field into 18 sub-regions, giving a total of $N_j = 54$ samples for each mock. The covariance matrix for an individual mock sample is estimated using equation (25), with an extra factor $(N_j - 1)$ to account for correlations between the Jackknife samples. We average over the covariance matrices of the 25 mocks to obtain the final covariance matrix. It is pointed out in Escoffier et al. (2016) that this method can reduce the noise on the covariance estimation, and fast approach the truth. However, we caution that these mocks are not completely independent, because they are constructed from the same N -body simulation (Gonzalez-Perez et al. 2014).

The posterior of the model parameters θ given data D is estimated in a Bayesian way:

$$P(\theta|D) = \frac{P(D|\theta)P(\theta)}{P(D)}, \quad (26)$$

where $P(\theta)$ is the prior and $P(D)$ is treated as a normalization. The term $P(D|\theta)$ is proportional to the likelihood $\mathcal{L}$, which we assume to be Gaussian:

$$\mathcal{L} \propto \exp(-\chi^2/2).$$

We use Monte Carlo Markov Chain (MCMC) sampling of the parameter space, implementing the python package *emcee*.⁵

4.3.2 Data compression

Instead of fitting the whole 2D correlation function, which requires an $N(r_p) \times N(\pi)$ – dimensional covariance matrix, we compress the 2D information into the multipoles defined as

$$\xi_\ell(r) = \frac{2\ell + 1}{2} \int_{-1}^1 \xi_c(r, \mu) \mathcal{P}_\ell(\mu) d\mu; \ell = 0, 2, 4. \quad (27)$$

We ignore higher order multipoles because they are typically noisy and more sensitive to non-linearity. Multipoles are computed by interpolating the 2D correlation function, and this is done consistently for both the measurements and the models. In the QD model, we exclude ξ_4 from the fitting, because the non-zero signal at scales $r \geq 10 h^{-1}$ Mpc cannot be well reproduced by this model.

We also considered adding the projected correlation function w_p :

$$w_p(r_p) = \int_{-\pi_{\max}}^{\pi_{\max}} \xi(r_p, \pi) d\pi. \quad (28)$$

This has the merit that it is in principle independent of RSD for large enough $\pi_{\max}$, and a knowledge of the true real-space clustering should be advantageous if we are focusing on redshift-space effects that cause deviations from this. However, we found in practice that

⁵<http://dfm.io/emcee/>

it was not possible to choose a large enough $\pi_{\max}$ to achieve results that converged to the true real-space clustering without the results being too noisy to be useful. The w_p statistic may be useful at small separations, $r < 10 h^{-1}$ Mpc, as a means of probing the real-space 1-halo term, but as discussed above we do not need to do this in the present work. We included w_p in the fitting for the QD model; however, due to the limited information it could provide in addition to the multipoles, the w_p was excluded in the HS model fitting.

For each of the six cross-correlation configurations, we fit the measurement simultaneously with its corresponding galaxy autocorrelation. This allows us to break the degeneracy between the galaxy and group bias.

4.3.3 Scale cuts

Below quasi-linear scales $r \sim 10 h^{-1}$ Mpc, both models may fail to capture the full non-linear features. Fitting data points at these scales may introduce significant bias into the measured growth rate. Therefore, we test the models on a set of minimum fitting scales, $r_{\min} = 2, 5, 10, 15, 20 h^{-1}$ Mpc using the mock catalogues, and adopt the most appropriate cut for each subsample. For the HS model, because the model is designed to be able to fit smaller scales, we only test the model at $r_{\min} = 2, 5, 10 h^{-1}$ Mpc.

4.3.4 Integral constraints

To account for the missing power for modes larger than the GAMA survey scale, we include the integral constraint I , which is a small constant added to the 2D correlation function. The expected integral constraint is given by

$$I \equiv \frac{b_1 b_2}{3} \sigma_{\text{eff}}^2 = \frac{b_1 b_2}{3} \int \Delta^2(k) W^2(k, r = r_{\text{eff}}) d \ln k, \quad (29)$$

where $\Delta^2(k)$ is the dimensionless linear matter power spectrum, b_1 and b_2 are the tracer biases, $W(k, r) = 3[\sin(kr)/(kr)^3 - \cos(kr)/(kr)^2]$, and r_{eff} is the effective radius of the survey volume in one of the GAMA fields, $V = 4\pi r_{\text{eff}}^3/3$. The factor 1/3 accounts for the fact that we combine measurements over three GAMA fields. Equation (29) gives $I = 0.0017b_1b_2$. This can also be measured in the mock data directly, by comparing the projected correlation function w_p at large scales between the average of 25 samples and the combination of all mock samples. The measured values are consistent with the expectation given the statistical errors. In the QD model, we allowed I to be a free parameter, and found that it has little impact on other parameters, with a posterior consistent with zero (e.g. see Tables D1 and D2 in Appendix D). The integral constraints are then fixed to the measured values from mock for the streaming model, as shown in Tables D3 and D4.

4.3.5 Priors

The parameters used in the two models and their uniform prior ranges can be found in Table 2. For each of the group-galaxy subsample, the galaxy autocorrelation is fitted simultaneously with the cross-correlation. Parameters with subscript ‘a’ are used for autocorrelations and ‘c’ for cross-correlations. There is another cosmological parameter that should be considered: the normalization of the (linear) matter power spectrum σ_8 . From equation (14), it is clear that on linear scales, σ_8 and b are completely degenerate, hence RSD measurements are usually quoted in the combination $f\sigma_8$. At large k , the shape of the non-linear power spectrum is actually

Table 2. Range of the uniform priors of the RSD fitting parameters. For growth rate, the usual constraint from RSD is $f\sigma_8$, but we fix $\sigma_8 = 0.81$ in this analysis. The α , and η parameters are the 1-halo parameters applied in the GSM model only.

Parameter	Prior (QD)	Prior (HS)
b_{gal}	[0.1, 2.5]	[0.5, 2]
b_{12}	[0.1, 2.5]	[0.25, 3]
f	[0, 2]	[0, 2]
σ_a (km s $^{-1}$)	[143, 1140]	[30, 800]
σ_c (km s $^{-1}$)	[143, 1140]	[30, 800]
I_a	[0, 0.1]	Fixed
I_c	[0, 0.1]	Fixed
α_a	–	[0.1, 2]
α_c	–	[0.1, 2]
η_a	–	[0.5, 2.5]
η_c	–	[0.5, 2.5]

sensitive to σ_8 . However, such dependence is weak for the scales probed here, and we fix $\sigma_8 = 0.81$ throughout the analysis.

5 RESULTS

5.1 Mocks

We fit both models to each of the 25 mock samples, and compute the mean and scatter of the best-fitting parameters. The aim is to assess the scale at which an unbiased growth rate can be recovered. The result is shown in Fig. 7 for the set of $r_{\min}$ as mentioned in previous sections, and for each of the six configurations. The fiducial value of f with ± 10 per cent range is marked by the grey band in each panel. The error bar is comparable to, but should not be taken directly as the expected error size on the GAMA sample. The specific values of all model parameters are summarized in Tables D1 and D3 in Appendix D.

Notice that in the case of halo streaming model, there is a caveat that the 1-halo templates are obtained from the average of the same set of mocks as they are tested on. Ideally, we would like to have access to multiple sets of simulations covering different cosmology and HOD prescription, with matched survey configurations as GAMA. Then, we would test the halo streaming model on by extracting the 1-halo templates from one set of simulations and apply it to the measurements from the others. In this way, we can assess whether the model is robust against bias due to a different cosmology or change of the simulation settings. Such test will be particularly relevant for the forthcoming large data sets, where the demand of the precision of the model is high. However, this is beyond the scope of this paper given the noise level of the GAMA data. We would like to defer such detailed comparison to a future study.

The top panel of Fig. 7 shows the recovered growth rate f using the QD model (Section 4.1). As expected, when the small scales are included ($r_{\min} \leq 5 h^{-1}$ Mpc), the fitted growth rates are significantly biased in all configurations, while at larger scales ($r_{\min} \geq 15 h^{-1}$ Mpc), they converge to the fiducial value. The overall growth rate seems to be underestimated by about 5–10 per cent for most scale cuts, but this is much smaller compared to the statistical error of the GAMA sample. It is noticeable that the blue configurations are less biased down to smaller scales, with f recovered to within 10 per cent at $r_{\min} = 5–10 h^{-1}$ Mpc, compared to the red configurations which are only unbiased at $r_{\min} = 15–20 h^{-1}$ Mpc. This may be due to the smaller FoG effect in the blue configurations compared to the red. From this test, we choose to adopt

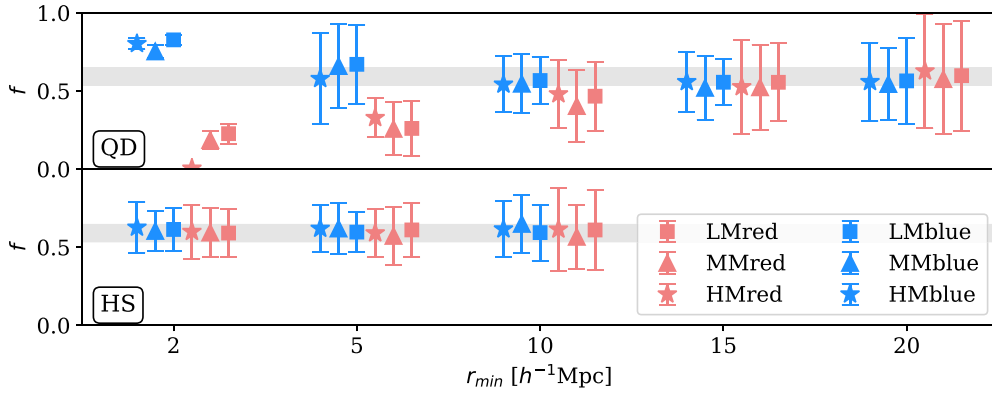


Figure 7. The means and scatter of the best-fitting growth rate f from 25 mocks as a function of the minimum fitting scale, $r_{\min}$, for the quasi-linear dispersion (QD: top) and the halo streaming (HS: bottom) models. Data points at each $r_{\min}$ are displaced by $0.1 h^{-1} \text{ Mpc}$ for clarity. The grey band marks the 10 per cent regions around the mock fiducial value $f = 0.593$ at $z = 0.195$. Note that the error bars are for a single survey, so that the errors on the mean of the mocks are 5 times smaller than shown.

$r_{\min} = \{10, 10, 10, 15, 20, 20\} h^{-1} \text{ Mpc}$ for the LMblue, MMblue, HMblue, LMred, MMred, and HMred subsamples, respectively, for the QD model in the application to the GAMA data. The bottom panel of Fig. 7 shows the recovered growth rate f using the HS model, where the results are impressively consistent. The growth rates for the different subsets are consistent to within an rms of 3 per cent in the mock average results, and the global average of these different subsets is within 2 per cent of the fiducial value. This successful performance holds down to even $r_{\min} = 2 h^{-1} \text{ Mpc}$, although the estimated errors at that point are little different to those at $r_{\min} = 5 h^{-1} \text{ Mpc}$, so we conservatively adopt the larger figure in our HS analysis.

Fig. 8 shows the linear and streaming model with the mean best-fitting parameters from the mock subsamples, at the respective $r_{\min}$ as mentioned above. The mock average measurement as well as the 1σ error on the mean is also shown. In addition, we also show the corresponding 2-halo term of the streaming model in dotted black lines. On large scales ($r \geq 15 h^{-1} \text{ Mpc}$), all of the model curves converge, and match well with the mock average. It is noticeable that the full streaming model (with the addition of the 1-halo template) and its 2-halo term do not coincide exactly on these scales: the extracted 1-halo template still has some residuals in the monopole and quadrupole. The largest difference is seen in the hexadecapole. The slightly positive values seem to be produced only by the 1-halo FoG, which both the linear and the 2-halo terms of the streaming model fail to capture. Looking at smaller scales ($r \leq 10 h^{-1} \text{ Mpc}$), it seems that the QD model underpredicts the power in the red configurations, and overpredicts that in the blue configurations.

5.2 GAMA

Fig. 9 shows the measured GAMA multipoles (filled circles), the best-fitting QD models (black dashed lines), and the HS models (black solid lines). In addition, the corresponding HS model 2-halo term is shown in the dotted black lines. The same scale cuts, $r_{\min}$, are adopted as in the mock case for each of the models. The χ^2 and parameter values are shown in Tables D2 and D4. The full HS model 2D models are contrasted with the GAMA data in Appendix C.

We see that the QD model provides a reasonable fit to the monopole and quadrupole at given $r_{\min}$ in most configurations. The only exception is the LMred and LMblue subsamples, where the monopole power is boosted at large scales and the quadrupole power

is consistent with zero. Despite the visual discrepancy, the χ^2 of these models are consistent with the degrees of freedom of the data. The HS model is able to capture the shape of the multipoles down to smaller scales, especially the hexadecapole at scales $r > 5 h^{-1} \text{ Mpc}$. At smaller scales ($r < 5 h^{-1} \text{ Mpc}$), although excluded from fitting, the mock 1-halo template continues to provide a reasonable fit to the red configurations. But this is not the case for the blue configurations, where the non-linear velocity dispersion seems to be stronger in the actual data compared to the mock catalogues. One possible explanation could be the impact of redshift measuring errors. These are not included in the mocks, and so any measured velocity dispersion in the real data only will include the redshift error in quadrature. The typical GAMA error is 50 km s^{-1} , but in detail Liske et al. (2015) showed that redshift errors can depend on spectral and target properties. The redshift error for galaxies classified as the ‘absorption’ type (i.e. the spectrum is dominated by absorption features) is 101 km s^{-1} , compared to the ‘emission’ type, which is 33 km s^{-1} . But red galaxies have a larger measured velocity dispersion, so the impact of redshift errors on the total measured dispersion will be less in that case.

Fig. 10 shows the mean and 1σ error on the model parameters from the MCMC posterior for the GAMA data, fitted at respective $r_{\min}$. The open and filled symbols denote parameter constraints from the QD model and the HS model, respectively. In the latter case, we also show the constraints measured using the 1-halo template from the ‘contaminated’ red galaxy sample (purple filled symbols). All sets of constraints show good consistency. Notice that the size of the error bar in the blue configurations is similar in both models, although the QD model has a scale cut at $10 h^{-1} \text{ Mpc}$ while the HS model at $5 h^{-1} \text{ Mpc}$. This is the consequence of the additional nuisance parameters added in the latter model. The specific parameter values, including the 1-halo parameters in the HS model, can be found in Tables D2 and D4. In Figs E1–E4 in Appendix E, we further show the full posteriors from MCMC for all parameters in both models, grouped by the red and blue configurations. In the HS model, the 1-halo parameters α and η have no primary degeneracy with the growth rate, although the growth rate can be shifted slightly through their small degeneracy with the velocity dispersion parameters. In practice, one would always marginalize over the 1-halo parameters.

The middle two panels show the measured group and galaxy biases in both models. The LM, MM, and HM group biases measured consistently between the red and blue configurations in both models.

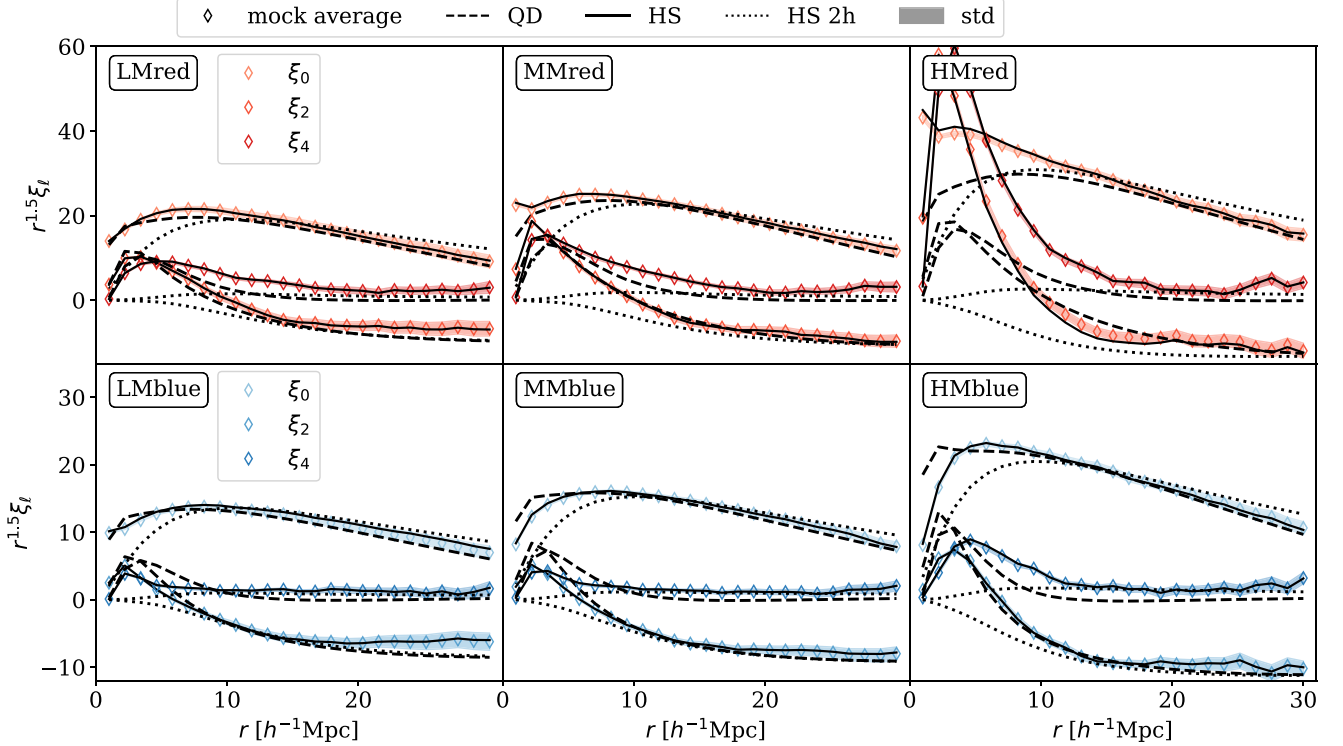


Figure 8. The multipoles of the group-galaxy cross-correlation functions in the mock average (open diamond). The coloured bands show the scatter on the mean. The best-fitting QD models are shown in dashed black line, with $r_{\min} = \{10, 10, 10, 15, 20, 20\} h^{-1} \text{Mpc}$ for the LMblue, MMblue, HMblue, LMred, MMred, and HMred subsamples, respectively. The best-fitting Gaussian streaming models with a 1-halo template are shown in solid black lines, with a fixed $r_{\min} = 5 h^{-1} \text{Mpc}$ for all sub-samples, with the corresponding 2-halo term shown in dotted black lines. For the presentation purpose, the multipoles have been multiplied by $r^{1.5}$.

For our selection of galaxies, we find that $b_{\text{gal}} \approx 1$ for the blue galaxies, and $b_{\text{gal}} \approx 1.3$ for the red galaxies; for the groups, we find that $b_{\text{grp}} \approx 1.2, 1.4, 1.8$ for the low, medium, and high mass ranges. The b_{grp} measurements are in qualitative agreement with that in Riggs et al. (2021) on large scales (e.g. see their Fig. 8), although direct comparison is non-trivial due to difference in the group selection. The consistency between the two models is good in general, although for the blue configuration, the HS model gives systematically lower biases compared to the QD model by $0.5\sigma - 1.5\sigma$. The lower two panels show the measured autocorrelations and cross-correlation velocity dispersion, σ_a and σ_c . The two models measure consistent velocity dispersion, despite slightly different form for the FoG term. Notice that for the red configuration in the QD model, because of the large scale cut, the velocity dispersion posterior is prior driven. There is a tentative trend (at $\sim 2\sigma$) that the red configurations have larger velocity dispersion compared to the blue configurations, with $\sigma_{a,c} \sim 400 - 500 \text{ km s}^{-1}$ for the red configurations, and $\sigma_{a,c} \sim 300 \text{ km s}^{-1}$ for the blue configurations in both autocorrelations and cross-correlations. There is, however, no clear dependence on the group mass. These measured galaxy biases and velocity dispersion are in good agreement with other measurements from GAMA (e.g. Blake et al. 2013; Loveday et al. 2018). The marginalized posterior for f , b_{gal} , and b_{12} for the six configurations is shown in Figs 11 and 12 for the QD and HS models, respectively.

The top panel shows the measured growth rate, consistent across the six subsamples for both models. Here, we have presented the results in the more general form of $f\sigma_8(z)$. The rationale for this is that

our modelling assumes that the background cosmology (WMAP7 parameters) is known exactly. This is not precisely true, and the observed distortion parameter, $\beta = f/b$, is actually $\propto f\sigma_8$ (since $b\sigma_8$ is observable). We therefore multiply our fitted f by the fiducial $\sigma_8(z)$ in order to obtain a combination that should be insensitive to the exact fiducial model.

We also note that RSD analyses commonly also allow for the Alcock-Paczynski effect (Alcock & Paczynski 1979; Ballinger, Peacock & Heavens 1996), which introduces additional distortions of the 2D correlation function from distance measurements using a ‘wrong’ cosmology. This degree of freedom can boost the errors on $f\sigma_8$ substantially if the cosmological model is left free. But the AP effect is unimportant if the model is constrained by precise external CMB data as here. A further reason that this is reasonable is that the interest in RSD comes from the desire to test gravity: the CMB data give a precise prediction of $f\sigma_8$ and we want to know if this is what we measure.

In detail, then, we take $\sigma_8(z) = g(z)\sigma_8(0) = 0.73$, where $z = 0.20$, $\sigma_8(0) = 0.81$, and $g(z)$ is the time-dependence of the (linear) density fluctuation in linear theory, normalized to $g(0) = 1$. The measurements give a mean of $f\sigma_8 = 0.27$ with uncertainties ranging from 0.07 to 0.20 for the QD model, and $f\sigma_8 = 0.29$ with uncertainties ranging from 0.07 to 0.14 for the HS model. We combine our measurements from the six cross-correlation configurations for the HS model, accounting for their correlations using the mock catalogues. We compute the scatter on the average as well as the covariance of the best-fitting f for the six configurations in the 25 mock samples. Ideally, one would like to use the full posterior. However, this would

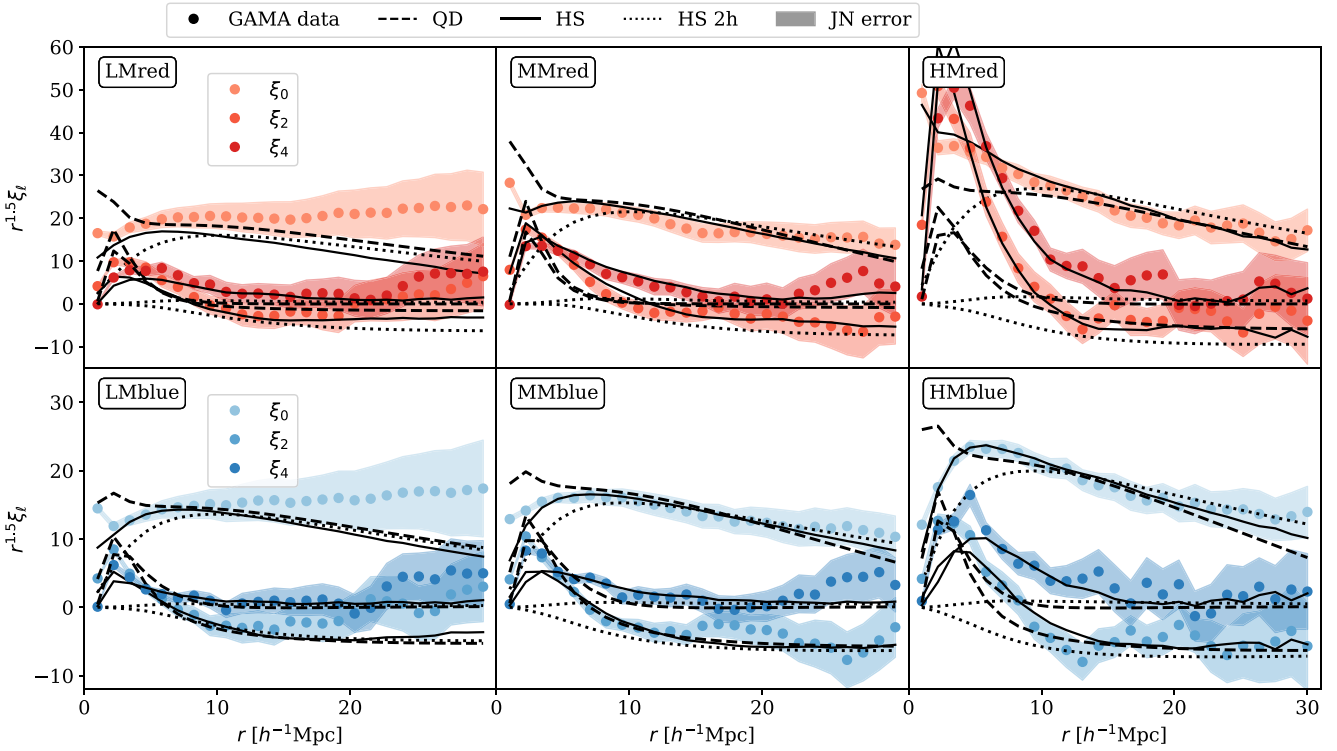


Figure 9. Same as Fig. 8 but for the actual GAMA data, with the same $r_{\min}$ adopted. The coloured bands show Jackknife errors. The χ^2 for each of the models can be found in Tables D2 and D4.

require the time consuming step of running MCMC for each of the mock sample, thus the simple average of the maximum likelihood values is adopted. Our combined constraint from the HS model thus gives

$$f\sigma_8(z = 0.20) = 0.29 \pm 0.10. \quad (30)$$

The corresponding figure for the QD model is 0.27 ± 0.14 , showing the extra information gained through the smaller scales that the HS model is able to probe. We note that the limited number of mock samples means that our covariance matrices will be imprecise, so that the errors on the growth rate for an individual sample may be underestimated (Hartlap, Simon & Schneider 2007; Sellentin & Heavens 2016). However, the empirical dispersion in the mean of the maximum-likelihood values should be robust.

The striking thing about this GAMA-based figure is that it is rather low compared to the fiducial *Planck* figure of $f\sigma_8(z = 0.20) = 0.47 \pm 0.01$, derived from the *Planck* TT, TE, EE+lowE+lensing cosmological parameters (Planck Collaboration VI 2020): our figure is 1.8σ below this *Planck* value. This discrepancy is certainly unexpected given how well our modelling was able to account for the RSD signal in the different mock realizations, and how the recovered growth rates were consistent between different methods of model fitting. Furthermore, the figures recovered from the different GAMA subsamples show the same level of consistency with each other as is seen in subsamples within the mocks.

There are a number of things that can be said about the low observed figure. The first is that there is some evidence that the fiducial *Planck* figure may be too high, with local gravitational lensing data consistently arguing for a reduction of about 10 per cent (see e.g. Hang et al. 2021). Our measurement would then be in 1.3σ disagreement with a revised fiducial value of 0.42, implying that GAMA is an unusual data set, but not unreasonably so. And we do

have evidence that this is the case: inspection of Fig. 3 shows that $N(z)$ has a substantial dip at $z \simeq 0.24$, which is seen consistently in all three fields. One might suspect a problem with the redshift pipeline, but this feature is absent in a subsequent fourth GAMA field, not used here; the three main GAMA fields are simply rather unusual regions of space. Finally, note that a multitracer analysis of RSD in GAMA by Blake et al. (2013) gave $f\sigma_8(z = 0.18) = 0.36 \pm 0.09$, which is also slightly lower than *Planck*, albeit not inconsistently so.

5.3 Group bias

Finally, it is interesting to ask if the group biases that we measure are in accord with what is expected for haloes of these masses. We compute the expected group bias in GAMA from the calibrated halo mass for the groups based on equation (6). We adopt the Tinker et al. (2010) fitting formula for the linear halo bias. The halo bias is expressed in terms of the peak height parameter $\nu \equiv \delta_c/\sigma_R$, where $\delta_c \approx 1.686$, and σ_R is the rms of the linear power spectrum filtered with a spherical top hat function with radius R (cf. equation 29). This is related to the halo mass via $M_h = 200\bar{\rho}_m 4\pi R^3/3$, where $\bar{\rho}_m$ is the mean background density. The mean group bias in a stellar mass range is estimated by

$$\hat{b}_{\text{grp}} = \frac{\sum_i N(\log M_h^i) b(\log M_h^i)}{\sum_i N(\log M_h^i)}, \quad (31)$$

where $N(\log M_h^i)$ is the number of groups in the logarithmic halo mass bin i . The halo mass adopted in case of mock and GAMA are as shown in Fig. 2. For GAMA, we also include the uncertainty in the calibrated halo mass due to the uncertainties in $\log M_p$ and α in equation (6). We combine the error via:

$$\Delta_{\log M_h} = \Delta_{\log M_p} + |\Delta_{\alpha} \log(L_{\text{grp}}/L_0)|. \quad (32)$$

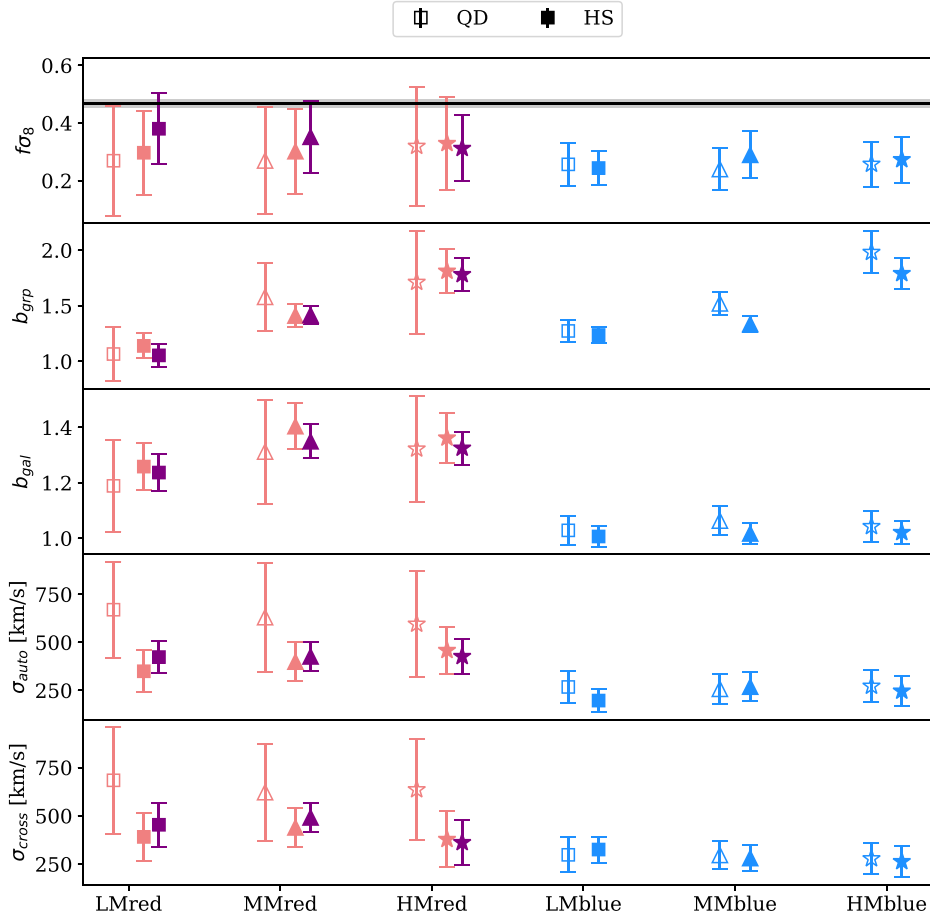


Figure 10. The MCMC parameter constraints for the actual GAMA data using the QD model (open symbols) and the streaming model with a 1-halo template (filled symbols). Each model is fitted with the respective $r_{\min}$ as in Fig. 8. The filled purple points show the constraints obtained using the 1-halo template from the ‘contaminated’ red sample in the mock data. The 1-halo parameters are marginalized over in the HS model. On the top panel, we have converted the constraints back to $f\sigma_8$ by multiplying back the fiducial $\sigma_8(z = 0.195)$. The black line and the grey band on the top panel mark $f\sigma_8 = 0.47 \pm 0.01$, the fiducial growth rate at $z = 0.195$ using Ω_m and σ_8 constraints from Planck Collaboration VI (2020). The specific values for all model parameters are shown in Tables D2 and D4.

For the mass range concerned here, $\Delta_{\log M_h} = 0.8\text{--}0.2$ dex from low mass to high mass. We account for this scatter by convolving the number of objects with a Gaussian distribution with width $\Delta_{\log M_h}$ in each $\log M_h$ bin. The predicted and measured group biases using mocks and GAMA data are summarized in Table 3. We also show biases computed at the logarithmic mean halo mass $\bar{M}_h$. These values are close to the average bias computed from equation (31) in the case of GAMA, but they deviate from the mock average significantly. This indicates that estimating the bias from the mean halo mass depends heavily on the distribution of the halo mass of the sample considered. In addition, we include the case where the more up-to-date mass–luminosity relation from Rana et al. (2022) is used to compute the group halo mass. The equivalent parameters to equation (6) are $M_p = (8.1 \pm 0.4) \times 10^{13} h^{-1} M_\odot$, $L_0 = 10^{11.5} h^{-2} L_\odot$, and $\alpha = 1.01 \pm 0.07$. The halo mass computed with this calibration is larger than using the fiducial Han et al. (2015), resulting in good consistency with the fitted group bias from both the QD and HS models as shown in Table 3.

In the mock catalogues, the predicted group bias for the LM, MM, and HM subsamples are all lower than the fitted values. These differences are significant given the error bar on the average of the mock measurements. In all cases, the group bias has apparently been

underestimated by about 20 per cent. This deviation in the group bias may arise because the arithmetic mean host halo mass of the group members is used as a proxy for the group halo mass. However, if one uses the total mass of unique host haloes in the group as M_h , then the bias in each mass range only increases by ~ 10 per cent. Since the mock group masses are calibrated using ‘real’ simulation halo masses, the difference illustrates clearly that our galaxy groups are not in 1-to-1 correspondence with single haloes, emphasizing once again the importance of analysing real and mock data with the same group finder.

6 SUMMARY AND CONCLUSIONS

In this work, we have investigated the RSD of group-galaxy cross-correlations, with the aim of understanding the robustness with which measurements of the density fluctuation growth rate can be extracted from such measurements. We have focused on the differences in the measured RSD using different types of galaxy and group, and developed new methods for fitting such data down to the small-scale non-linear regime.

We have used data from the GAMA survey in the redshift range $0.1 < z < 0.3$ to measure the 2D cross-correlation function $\xi(r_p)$,

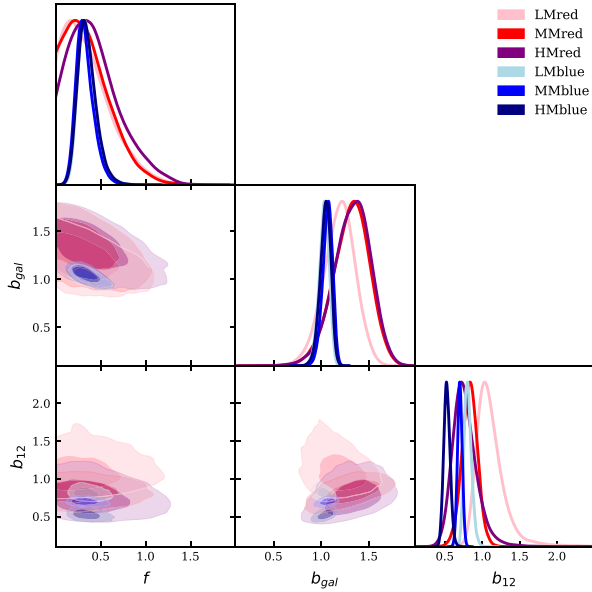


Figure 11. Marginalized MCMC posteriors for the QD model.

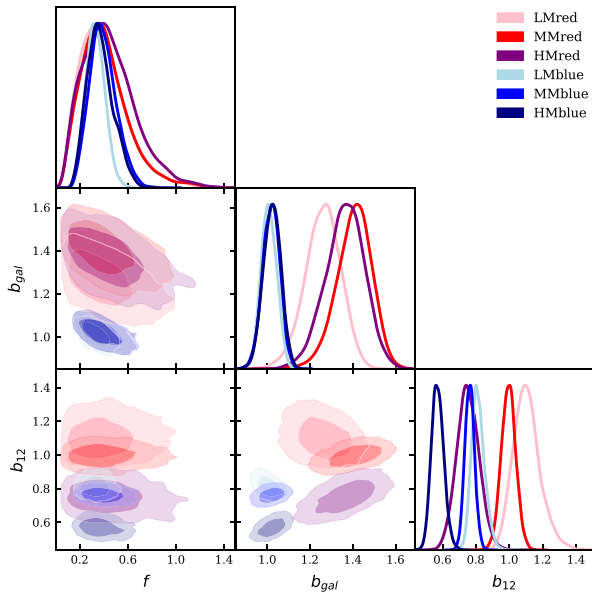


Figure 12. Marginalized MCMC posteriors for the HS model. The best-fitting bias parameters are very different among different sub-samples, but the recovered f values are consistent.

π) between groups and galaxies. The groups were found using an FoF algorithm from Treyer et al. (2018), and were subdivided into three stellar mass bins (LM: 40 per cent, MM: 50 per cent, and HM: 10 per cent). The corresponding halo mass for the groups was calibrated using the relation in Han et al. (2015), and the groups are expected to have typical masses of ($10^{12.2}$, $10^{12.7}$, $10^{13.2}$) $M_{\odot}$. For Rana et al. (2022), the mean halo masses are ($10^{12.6}$, $10^{13.1}$, $10^{13.5}$) $M_{\odot}$. The galaxies were split into red and blue subsets using a cut in the $g - i$ versus z plane, yielding in total six cross-correlation configurations: LMred, MMred, HMred, LMblue, MMblue, and HMblue.

We have used 25 GAMA light-cone mocks from Farrow et al. (2015) to test RSD models and to construct Jackknife covariance matrices for likelihood fitting. Mock group catalogues were generated using the identical algorithm that was applied to the real GAMA data. The mock catalogues are distinct from observation in several aspects: the mean redshift distribution, the bimodal $g - i$ colour distribution, and the total stellar mass of the groups. We discuss the appropriate empirical selection that yields the best match between the mocks and the data subsamples. The measured 2D correlation functions show good consistency between the data and the mocks down to small scales, and the same variation of the signals with galaxy colours and group masses are observed. The different cross-correlation results yield group biases that increase with mass, as expected. For GAMA, the predicted group bias from Tinker et al. (2010) is lower but consistent with the fitted values using the halo mass calibration in Han et al. (2015), whereas that from Rana et al. (2022) agree well with the fitted values. For mocks, however, these values tend to be higher than predicted. This difference illustrates that the groups found in redshift space do not constitute a pure halo sample.

We have compared these measurements with two RSD models: (1) a quasi-linear dispersion (QD) model; (2) a novel halo streaming (HS) model. The QD model is a generalization of the linear dispersion model of Mohammad et al. (2016) to use the non-linear real-space power spectrum. We found from testing on the mocks that this model provides unbiased measurements of the growth rate at $r_{\min} = 10\text{--}20 h^{-1}$ Mpc depending on the subsample. The HD model uses a halo model decomposition of the correlations, where a streaming model 2-halo term is combined with an empirical 1-halo template adopted from the mock average. This promising model, with the addition of two nuisance parameters, allows unbiased results on the growth rate down to $r_{\min} = 5 h^{-1}$ Mpc when fitting individual mock realizations, and for all group-galaxy combinations. For the GAMA measurements, an MCMC analysis was used to obtain the posterior of our model parameters. We found that given the scale cuts, all of the subsamples recover consistent growth rates in both models. The average growth rate from the six subsamples using the HS model is $f\sigma_8 = 0.29 \pm 0.10$ at $z = 0.20$, where the error should be robust as it is taken directly from the dispersion in maximum-likelihood values for the mock data. This figure is 1.8σ lower than the fiducial *Planck* value of $f\sigma_8 = 0.47 \pm 0.01$, and we have considered the implications of this result. At face value, the low GAMA result is consistent with the suggestions from gravitational lensing that the true value of $f\sigma_8$ may be about 10 per cent lower than the *Planck* central figure (e.g. Hang et al. 2021). But there are objective reasons to believe that the GAMA data set may be a statistical outlier, based on known anomalies in the redshift distribution in the GAMA fields.

Therefore, the real test of the RSD modelling presented here will be when it can be applied to much larger and more precise data sets, such as the Bright Galaxy Sample from the Dark Energy Spectroscopic Instrument (DESI) survey (Martini et al. 2018) and the Wide Area VISTA Extra-Galactic Survey (WAVES; Driver et al. 2016). We are greatly encouraged by the success of our halo streaming model in reproducing mock cross-correlations down to the smallest scales, and in yielding consistent values of $f\sigma_8$ from different tracers, to a tolerance of better than 3 per cent. This hybrid approach, taking advantage of ever more realistic mock data, therefore seems an attractive way of obtaining robust constraints on the growth of cosmological density fluctuations, and we look forward to seeing it applied to next-generation surveys.

Table 3. Bias for the groups in the LM, MM, and HM stellar mass bins for the mock average and GAMA. The first column shows the mean bias value computed from the fitting formula (Tinker et al. 2010), and the second column shows the corresponding bias computed at the mean halo mass in each case. The next two columns marked with a '*' show the bias computed in the same way, but with the GAMA group halo mass computed from a more up-to-date mass–luminosity relation (Rana et al. 2022). The rest of the columns show the fitted biases from the six cross-correlation configurations in the mocks and the GAMA data.

Group bias		T10	T10 $b(\bar{M}_h)$	T10*	T10* $b(\bar{M}_h)$	QD-red	QD-blue	HS-red	HS-blue
Mocks	LM	1.02	0.92	–	–	1.20 ± 0.04	1.20 ± 0.02	1.18 ± 0.03	1.20 ± 0.02
	MM	1.26	1.13	–	–	1.48 ± 0.05	1.46 ± 0.02	1.42 ± 0.02	1.38 ± 0.02
	HM	1.83	1.65	–	–	1.90 ± 0.06	2.09 ± 0.03	1.96 ± 0.03	1.92 ± 0.03
GAMA	LM	1.00	0.96	1.12	1.11	1.07 ± 0.24	1.27 ± 0.10	1.14 ± 0.11	1.24 ± 0.08
	MM	1.20	1.17	1.41	1.39	1.58 ± 0.30	1.52 ± 0.10	1.41 ± 0.10	1.34 ± 0.07
	HM	1.52	1.49	1.85	1.81	1.71 ± 0.46	1.98 ± 0.19	1.81 ± 0.20	1.79 ± 0.14

ACKNOWLEDGEMENTS

QH was supported by Edinburgh Global Research Scholarship and the Higgs Scholarship from Edinburgh University. JAP and SA were supported by the European Research Council under the COSFORM grant no. 670193. YC acknowledges the support of the Royal Society through a University Research Fellowship and an Enhancement Award. YC thanks the hospitality of the Astrophysics and Theoretical Physics groups of the Department of Physics at the Norwegian University of Science and Technology during his visit. MB is supported by the Polish National Science Center through grants no. 2020/38/E/ST9/00395, 2018/30/E/ST9/00698, 2018/31/G/ST9/03388, and 2020/39/B/ST9/03494, and by the Polish Ministry of Science and Higher Education through grant DIR/WK/2018/12.

GAMA is a joint European-Australasian project based around a spectroscopic campaign using the Anglo-Australian Telescope. The GAMA input catalogue is based on data taken from the *Sloan Digital Sky Survey* and the UKIRT Infrared Deep Sky Survey. Complementary imaging of the GAMA regions is being obtained by a number of independent survey programmes including GALEX MIS, VST KiDS, VISTA VIKING, *WISE*, *Herschel*-ATLAS, GMRT and ASKAP providing UV to radio coverage. GAMA is funded by the STFC (UK), the ARC (Australia), the AAO, and the participating institutions. The GAMA website is <http://www.gama-survey.org/>.

DATA AVAILABILITY

All of the GAMA data are publicly available on <http://www.gama-survey.org/dr3/schema/>. The mock catalogues and the GAMA group catalogue are available upon reasonable request.

REFERENCES

Alam S. et al., 2017, *MNRAS*, 470, 2617
 Alam S. et al., 2021a, *Phys. Rev. D*, 103, 083533
 Alam S., Peacock J. A., Farrow D. J., Loveday J., Hopkins A. M., 2021b, *MNRAS*, 503, 59
 Alcock C., Paczynski B., 1979, *Nature*, 281, 358
 Baldry I. K. et al., 2018, *MNRAS*, 474, 3875
 Ballinger W. E., Peacock J. A., Heavens A. F., 1996, *MNRAS*, 282, 877
 Beutler F., Dio E. D., 2020, *J. Cosmol. Astropart. Phys.*, 2020, 048
 Beutler F. et al., 2012, *MNRAS*, 423, 3430
 Bianchi D., Chiesa M., Guzzo L., 2015, *MNRAS*, 446, 75
 Bilicki M. et al., 2021, *A&A*, 653, A82
 Blake C. et al., 2011, *MNRAS*, 415, 2876
 Blake C. et al., 2013, *MNRAS*, 436, 3089
 Bond J. R., Kofman L., Pogogyan D., 1996, *Nature*, 380, 603
 Bonvin C., Hui L., Gaztañaga E., 2014, *Phys. Rev. D*, 89, 083535

Boylan-Kolchin M., Springel V., White S. D. M., Jenkins A., Lemson G., 2009, *MNRAS*, 398, 1150
 Cai Y.-C., Kaiser N., Cole S., Frenk C., 2017, *MNRAS*, 468, 1981
 Camacho H. et al., 2019, *MNRAS*, 487, 3870
 Carrasco J. J. M., Hertzberg M. P., Senatore L., 2012, *J. High Energy Phys.*, 2012, 82
 Cole S., 2011, *MNRAS*, 416, 739
 Cooray A., Sheth R., 2002, *Phys. Rep.*, 372, 1
 Cuesta-Lazaro C., Li B., Eggemeier A., Zarrouk P., Baugh C. M., Nishimichi T., Takada M., 2020, *MNRAS*, 498, 1175
 d'Amico G., Gleyzes J., Kokron N., Markovic K., Senatore L., Zhang P., Beutler F., Gil-Marín H., 2020, *J. Cosmol. Astropart. Phys.*, 2020, 005
 Davis M., Peebles P. J. E., 1983, *ApJ*, 267, 465
 Davis M., Nusser A., Masters K. L., Springob C., Huchra J. P., Lemson G., 2011, *MNRAS*, 413, 2906
 de la Torre S., Guzzo L., 2012, *MNRAS*, 427, 327
 Dekel A., Lahav O., 1999, *ApJ*, 520, 24
 Driver S. P. et al., 2011, *MNRAS*, 413, 971
 Driver S. P., Davies L. J., Meyer M., Power C., Robotham A. S. G., Baldry I. K., Liske J., Norberg P., 2016, in Napolitano N. R., Longo G., Marconi M., Paolillo M., Iodice E., eds. *Astrophysics and Space Science Proceedings Vol. 42, The Universe of Digital Sky Surveys*. Springer Cham, Switzerland, p. 205
 Driver S. P. et al., 2022, *MNRAS*, 513, 439
 Eke V. R. et al., 2004, *MNRAS*, 348, 866
 Escoffier S. et al., 2016, preprint ([arXiv:1606.00233](https://arxiv.org/abs/1606.00233))
 Farrow D. J. et al., 2015, *MNRAS*, 454, 2120
 Fisher K. B., 1995, *ApJ*, 448, 494
 Gonzalez-Perez V., Lacey C. G., Baugh C. M., Lagos C. D. P., Helly J., Campbell D. J. R., Mitchell P. D., 2014, *MNRAS*, 439, 264
 Guo H., Zehavi I., Zheng Z., 2012, *ApJ*, 756, 127
 Guo H. et al., 2014, *MNRAS*, 441, 2398
 Guzzo L. et al., 2008, *Nature*, 451, 541
 Hamilton A. J. S., 1992, *ApJ*, 385, L5
 Han J. et al., 2015, *MNRAS*, 446, 1356
 Hand N., Seljak U., Beutler F., Vlah Z., 2017, *J. Cosmol. Astropart. Phys.*, 2017, 009
 Hang Q., Alam S., Peacock J. A., Cai Y.-C., 2021, *MNRAS*, 501, 1481
 Hartlap J., Simon P., Schneider P., 2007, *A&A*, 464, 399
 Hawkins E. et al., 2003, *MNRAS*, 346, 78
 Jain B., Khoury J., 2010, *Ann. Phys.*, 325, 1479
 Kaiser N., 1987, *MNRAS*, 227, 1
 Komatsu E. et al., 2011, *ApJS*, 192, 18
 Kraljic K. et al., 2018, *MNRAS*, 474, 547
 Lahav O., Lilje P. B., Primack J. R., Rees M. J., 1991, *MNRAS*, 251, 128
 Landy S. D., Szalay A. S., 1993, *ApJ*, 412, 65
 Lemson G., Virgo Consortium t., 2006, preprint ([astro-ph/0608019](https://arxiv.org/abs/astro-ph/0608019))
 Linder E. V., 2005, *Phys. Rev. D*, 72, 043529
 Liske J. et al., 2015, *MNRAS*, 452, 2087
 Loveday J. et al., 2012, *MNRAS*, 420, 1239
 Loveday J. et al., 2018, *MNRAS*, 474, 3435
 Mandelbaum R., Wang W., Zu Y., White S., Henriques B., More S., 2016, *MNRAS*, 457, 3200

- Martini P. et al., 2018, in Evans C. J., Simard L., Takami H., eds, *Proc. SPIE Conf. Ser. Vol. 10702, Ground-based and Airborne Instrumentation for Astronomy VII*. SPIE, Bellingham, p. 107021F
- McDonald P., Seljak U., 2009, *J. Cosmol. Astropart. Phys.*, 10, 007
- Merson A. I. et al., 2013, *MNRAS*, 429, 556
- Mohammad F. G., de la Torre S., Bianchi D., Guzzo L., Peacock J. A., 2016, *MNRAS*, 458, 1948
- Padilla N. D., Merchán M. E., Valotto C. A., Lambas D. G., Maia M. A. G., 2001, *ApJ*, 554, 873
- Parkinson D. et al., 2012, *Phys. Rev. D*, 86, 103518
- Peacock J. A., 2016, in van de Weygaert R., Shandarin S., Saar E., Einasto J., eds, *Proc. IAU Symp. 308, The Zeldovich Universe: Genesis and Growth of the Cosmic Web*. Kluwer, Dordrecht, p. 125
- Peacock J. A., Dodds S. J., 1994, *MNRAS*, 267, 1020
- Peacock J. A., Smith R. E., 2000, *MNRAS*, 318, 1144
- Peacock J. A. et al., 2001, *Nature*, 410, 169
- Pezzotta A. et al., 2017, *A&A*, 604, A33
- Planck Collaboration VI, 2020, *A&A*, 641, A6
- Rana D., More S., Miyatake H., Nishimichi T., Takada M., Robotham A. S. G., Hopkins A. M., Holwerda B. W., 2022, *MNRAS*, 510, 5408
- Reid B. A., White M., 2011, *MNRAS*, 417, 1913
- Reid B. A. et al., 2012, *MNRAS*, 426, 2719
- Reid B. A., Seo H.-J., Leauthaud A., Tinker J. L., White M., 2014, *MNRAS*, 444, 476
- Riggs S. D., Barbhuiyan R. W. Y. M., Loveday J., Brough S., Holwerda B. W., Hopkins A. M., Phillipps S., 2021, *MNRAS*, 506, 21
- Robotham A. S. G. et al., 2011, *MNRAS*, 416, 2640
- Scoccimarro R., 2004, *Phys. Rev. D*, 70, 083007
- Seljak U., 2000, *MNRAS*, 318, 203
- Sellentin E., Heavens A. F., 2016, *MNRAS*, 456, L132
- Senatore L., 2015, *J. Cosmol. Astropart. Phys.*, 2015, 007
- Sheth R. K., 1996, *MNRAS*, 279, 1310
- Smith R. E. et al., 2003, *MNRAS*, 341, 1311
- Takahashi R., Sato M., Nishimichi T., Taruya A., Oguri M., 2012, *ApJ*, 761, 152
- Taruya A., Nishimichi T., Saito S., 2010, *Phys. Rev. D*, 82, 063522
- Taylor E. N. et al., 2011, *MNRAS*, 418, 1587
- Taylor E. N. et al., 2015, *MNRAS*, 446, 2144
- Tinker J. L., Robertson B. E., Kravtsov A. V., Klypin A., Warren M. S., Yepes G., Gottlöber S., 2010, *ApJ*, 724, 878
- Treyer M. et al., 2018, *MNRAS*, 477, 2684
- Viola M. et al., 2015, *MNRAS*, 452, 3529
- Vlah Z., Castorina E., White M., 2016, *J. Cosmol. Astropart. Phys.*, 2016, 007
- Wang L., Steinhardt P. J., 1998, *ApJ*, 508, 483
- Wechsler R. H., Tinker J. L., 2018, *ARA&A*, 56, 435
- Wojtak R., Hansen S. H., Hjorth J., 2011, *Nature*, 477, 567
- Yang X., Mo H. J., van den Bosch F. C., Weinmann S. M., Li C., Jing Y. P., 2005, *MNRAS*, 362, 711
- Zu Y., Weinberg D. H., 2013, *MNRAS*, 431, 3319

SUPPORTING INFORMATION

Supplementary data are available at [MNRAS](https://academic.oup.com/mnras/article/517/1/374/6696392) online.

Supplementary materials.pdf

Please note: Oxford University Press is not responsible for the content or functionality of any supporting materials supplied by the authors. Any queries (other than missing material) should be directed to the corresponding author for the article.

This paper has been typeset from a $\text{\LaTeX}$ file prepared by the author.